



THE INTERNET THAT NEVER WAS Why OSI remained a "beautiful dream" P. 36	FLYING HIGH ON LIGHT ALONE The age of the solar airplane is dawning P. 07	THE GREAT MAGNETIC DIVIDE Have we finally made a monopoly? P. 42	YOUR BIONIC FUTURE AWAITS Hear the exclusive podcast series at SPECTRUM.IEEE.ORG
--	---	--	--

IEEE
SPECTRUM

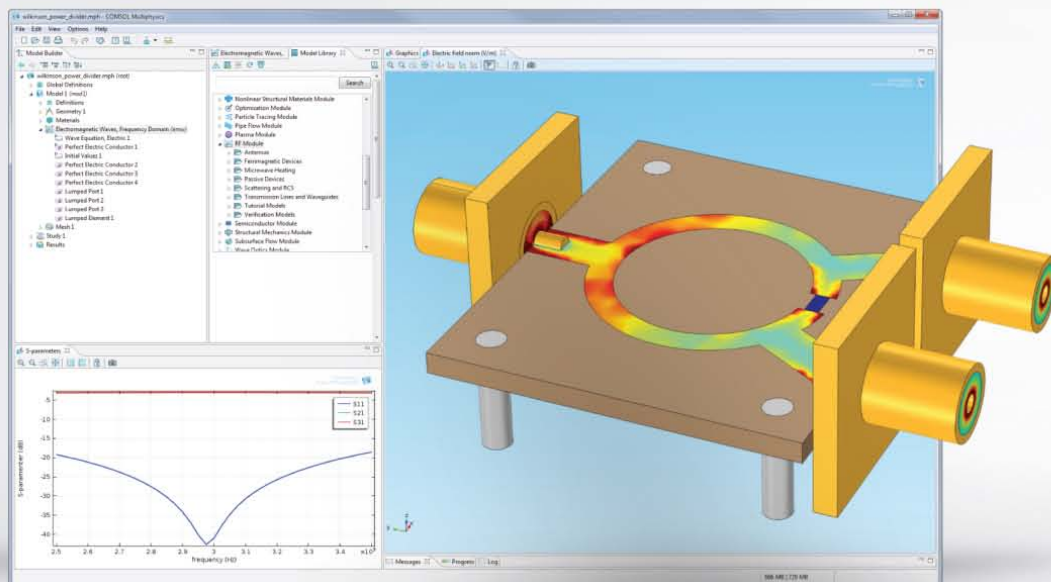
FOR THE TECHNOLOGY INSIDER | 08.13

GHOST SHIP
THE USS ZUMWALT
IS A NEW KIND
OF STEALTHY
ELECTRIC WARSHIP



COMSOL 4.3b
Now Available!
www.comsol.com/4.3b

RF DESIGN: Simulation results show the electric field distribution on top of the microstrip lines of a Wilkinson power divider. The S-parameters show input matching at 3 GHz and evenly divided power at the two output ports.



Verify and optimize your designs with COMSOL Multiphysics®

Multiphysics tools let you build simulations that accurately replicate the important characteristics of your designs. The key is the ability to include all physical effects that exist in the real world. To learn more about COMSOL Multiphysics, visit www.comsol.com/introvideo

Product Suite

COMSOL Multiphysics

ELECTRICAL

AC/DC Module
RF Module
Wave Optics Module
MEMS Module
Plasma Module
Semiconductor Module

MECHANICAL

Heat Transfer Module
Structural Mechanics Module
Nonlinear Structural Materials Module
Geomechanics Module
Fatigue Module
Multibody Dynamics Module
Acoustics Module

FLUID

CFD Module
Microfluidics Module
Subsurface Flow Module
Pipe Flow Module
Molecular Flow Module

CHEMICAL

Chemical Reaction Engineering Module
Batteries & Fuel Cells Module
Electrodeposition Module
Corrosion Module
Electrochemistry Module

MULTIPURPOSE

Optimization Module
Material Library
Particle Tracing Module

INTERFACING

LiveLink™ for MATLAB®
LiveLink™ for Excel®
CAD Import Module
ECAD Import Module
LiveLink™ for SolidWorks®
LiveLink™ for SpaceClaim®
LiveLink™ for Inventor®
LiveLink™ for AutoCAD®
LiveLink™ for Creo™ Parametric
LiveLink™ for Pro/ENGINEER®
LiveLink™ for Solid Edge®
File Import for CATIA® V5

 **COMSOL**

FEATURES_08.13

IEEE
SPECTRUM

30 CLAD IN CONTROVERSY

Take a peek at the first of a new and contentious class of high-tech warships, a destroyer named the *USS Zumwalt*.

BY DAVID SCHNEIDER

24 The Cognitive Net Is Coming

Engineers are turning to human biology to design a bigger, faster, and more adaptable Internet.

By Antonio Liotta

36 The Internet That Wasn't

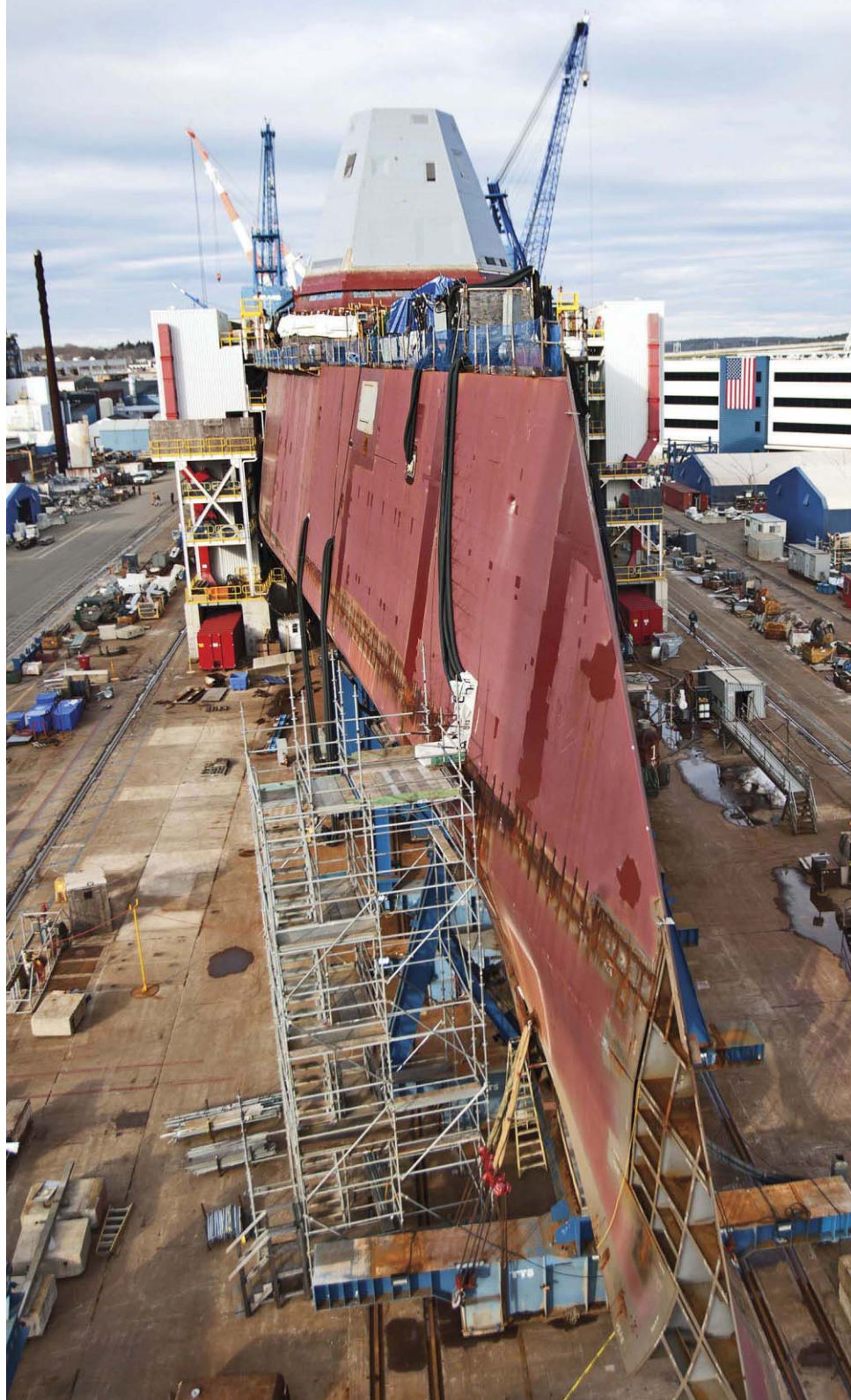
An international effort to set standards for computer networking had the support of everyone who counted—and that was the problem.

By Andrew L. Russell

42 The Race for the Pole

In some very special materials, north and south magnetic poles can move about independently; this could lead to entirely new classes of devices.

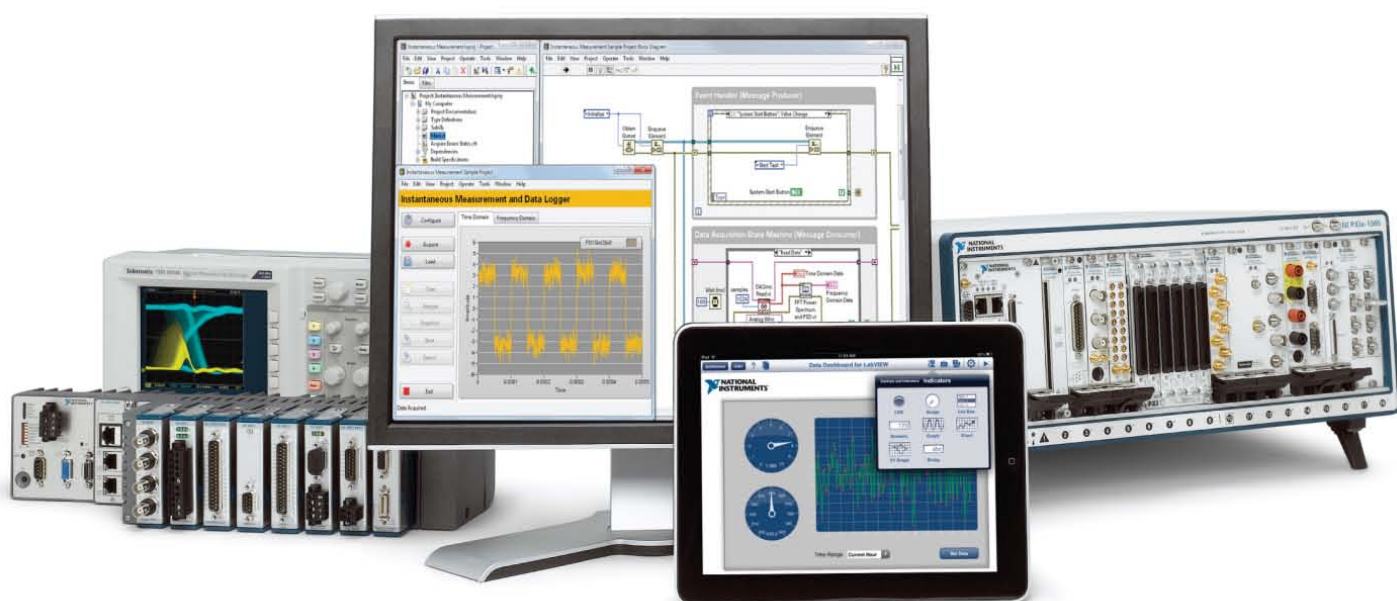
By Jonathan Morris



On the Cover Illustration for IEEE Spectrum by John MacNeill

Infinite Designs, One Platform

with the only complete
system design environment



NI LabVIEW is the only comprehensive development environment with the unprecedented hardware integration and wide-ranging compatibility you need to meet any measurement and control application challenge. And LabVIEW is at the heart of the graphical system design approach, which uses an open platform of productive software and reconfigurable hardware to accelerate the development of your system.

LabVIEW system design software offers unrivaled hardware integration and helps you program the way you think—graphically.



>> Accelerate your system design productivity at ni.com/labview-platform

800 453 6202

©2013 National Instruments. All rights reserved. LabVIEW, National Instruments, NI, and ni.com are trademarks of National Instruments. Other product and company names listed are trademarks or trade names of their respective companies. 11215

**NATIONAL
INSTRUMENTS™**

DEPARTMENTS_08.13

IEEE
SPECTRUM

07

News

Top Electric Airplanes

The continental crossing of Solar Impulse is just one indication of an electric plane revolution.

By Michele Traverso

- 09 Ocean Thermal Energy in China
- 10 Practical Quantum Photography
- 12 Intel Strikes Back
- 14 The Big Picture

17

Resources

Television on the Go

The rollout of an upgrade to digital TV brings broadcasts to smartphones.

By Stephen Cass

- 18 Hands On: Building a Beacon
- 20 At Work: Companies Adapt to "Bring Your Own Device"
- 21 Start-ups: Veebot
- 52 Dataflow: The Shape of Progress

06

Opinion

Spectral Lines

Two cheers for electric cars: Changing gears in the debate started by last month's cover story.

By John Voelcker

- 04 Back Story
- 05 Contributors
- 22 Technically Speaking

Online

Spectrum.ieee.org

The Making of Facebook's Graph Search

In 2011, Facebook CEO Mark Zuckerberg began dreaming of a natural-language search engine to explore the wealth of data posted on the site. Read how Facebook's engineers made that dream come true by developing Graph Search, which rolls out to U.S. users this month.

ADDITIONAL RESOURCES

Tech Insider / Webinars

Available at spectrum.ieee.org/webinar

- ▶ How an MBA Can Advance Your Engineering Career—1 August
- ▶ Modeling Batteries and Fuel Cells Using COMSOL Multiphysics—8 August
- ▶ Four Ways to Improve Linux Performance for Multicore Devices
- ▶ Delivering Full-Featured Mobile Experiences on Mid-Range Devices
- ▶ Integrated Distribution GIS and Enterprise Asset Management Proof of Concept at Southern Company
- ▶ Challenges and Solution Needs for Delivering the Promise of Next-Generation Infotainment
- ▶ **NEW PRODUCT RELEASE LIBRARY**
<http://spectrum.ieee.org/static/new-product-release-library>
- ▶ **MASTER BOND WHITE PAPER LIBRARY**
<http://spectrum.ieee.org/static/masterbond-whitepaper-library>

The Institute

Available 7 August at theinstitute.ieee.org

- ▶ **PROSTHETIC HAND MIMICS SENSE OF TOUCH** Neurobiologists at the University of Chicago successfully tested two types of sensors on a prosthetic hand that measure contact and pressure, which could soon allow people with prosthetics to be able to feel objects in their hands.
- ▶ **MENTOR CENTRE LAUNCHES** The new and improved IEEE mentoring program matches members who want to serve as mentors with students and young professionals who need career guidance. The system identifies those who have similar technical or volunteering interests and specific experience.
- ▶ **CLOUD COMPUTING IN NEW MARKETS** The IEEE International Conference on Cloud Computing in Emerging Markets, to be held from 16 to 18 October in Bangalore, India, will cover such topics as the design of cloud computing services, software and business process as a service, and advances in the virtualization of hardware and software services.

IEEE SPECTRUM

(ISSN 0018-9235) is published monthly by The Institute of Electrical and Electronics Engineers, Inc. All rights reserved. © 2013 by The Institute of Electrical and Electronics Engineers, Inc., 3 Park Avenue, New York, NY 10016-5997, U.S.A. Volume No. 50, issue No. 8, International edition. The editorial content of IEEE Spectrum magazine does not represent official positions of the IEEE or its organizational units. Canadian Post International Publications Mail (Canadian Distribution) Sales Agreement No. 40013087. Return undeliverable Canadian addresses to: Circulation Department, IEEE Spectrum, Box 1051, Fort Erie, ON L2A 6C7. Cable address: ITRIPLEE. Fax: +1 212 419 7570. INTERNET: spectrum@ieee.org. ANNUAL SUBSCRIPTIONS: IEEE Members: \$21.40 included in dues. Libraries/institutions: \$399. POSTMASTER: Please send address changes to IEEE Spectrum, c/o Coding Department, IEEE Service Center, 445 Hoes Lane, Box 1331, Piscataway, NJ 08855. Periodicals postage paid at New York, NY, and additional mailing offices. Canadian GST # 125634188. Printed at 120 Donnelly Dr., Glasgow, KY 42141-1060, U.S.A. IEEE Spectrum circulation is audited by BPA Worldwide. IEEE Spectrum is a member of the Association of Business Information & Media Companies, the Association of Magazine Media, and Association Media & Publishing. IEEE prohibits discrimination, harassment, and bullying. For more information, visit <http://www.ieee.org/web/aboutus/whatis/policies/p9-26.html>.

FROM LEFT: PIPISTREL; TAKA; GABRIEL TRIONE/

BACK STORY_



How Quickly We Forget

History is written by the winners, as they say. And in the fast-moving world of technology, history can mean things that happened just 15 or 20 years ago. In “The Internet That Wasn’t,” in this issue, Andrew L. Russell, an assistant professor of history and director of the Program in Science & Technology Studies at Stevens Institute of Technology, in Hoboken, N.J., explores just such a case: an alternative scheme for computer networking that, despite years of effort by thousands of engineers, ultimately lost out to the Internet’s Transmission Control Protocol/Internet Protocol (TCP/IP) and is now all but forgotten.

Russell first wrote about the competition between that scheme, called Open Systems Interconnection (OSI), and the Internet in 2006, for the *IEEE Annals of the History of Computing*. During his research on the Internet and its precursor, the ARPANET, “OSI would creep up as a foil, something they didn’t want the Internet to turn into,” he says. “So that’s the way I presented it.”

After the article was published, he says, veterans of OSI “came out of the woodwork to tell their stories.” One of the e-mails was from a computer networking pioneer named John Day, who had worked on both TCP/IP and OSI. Day told Russell that his article hadn’t captured the full scope of the story.

“Nobody likes to hear that they got it wrong,” Russell recalls. “It took me a while to cool down.” Eventually, he talked to Day, who put him in touch with other OSI participants in the United States and France. Through those interviews and archival research at the Charles Babbage Institute, in Minnesota, a more balanced, complex history of networking emerged, which he describes in his upcoming book *Open Standards and the Digital Age: History, Ideology, and Networks* (Cambridge University Press).

“It’s almost alarming that something that recent can be so easily forgotten,” Russell says. On the other hand, it’s what makes being a historian of technology so rewarding. ■

CITING ARTICLES IN IEEE SPECTRUM *IEEE Spectrum* publishes an international and a North American edition, as indicated at the bottom of each page. Both have the same editorial content, but because of differences in advertising, page numbers may differ. In citations, you should include the issue designation. For example, Dataflow is in *IEEE Spectrum*, Vol. 50, no. 8 (INT), August 2013, p. 52, or in *IEEE Spectrum*, Vol. 50, no. 8 (NA), August 2013, p. 56.

IEEE
SPECTRUM

EDITOR IN CHIEF

Susan Hassler, s.hassler@ieee.org

EXECUTIVE EDITOR

Glenn Zorpette, g.zorpette@ieee.org

EDITORIAL DIRECTOR, DIGITAL

Harry Goldstein, h.goldstein@ieee.org

MANAGING EDITOR

Elizabeth A. Bretz, e.bretz@ieee.org

SENIOR ART DIRECTOR

Mark Montgomery, m.montgomery@ieee.org

SENIOR EDITORS

Stephen Cass (Resources), cass.s@ieee.orgErico Guizzo (Digital), e.guizzo@ieee.orgJean Kumagai, j.kumagai@ieee.orgSamuel K. Moore (News), s.k.moore@ieee.orgTekla S. Perry, t.perry@ieee.orgJoshua J. Romero (Interactive), j.j.romero@ieee.orgPhilip E. Ross, p.ross@ieee.orgDavid Schneider, d.a.schneider@ieee.org

SENIOR ASSOCIATE EDITORS

Steven Cherry, s.cherry@ieee.org

DEPUTY ART DIRECTOR

Brandon Palacio

PHOTO & MULTIMEDIA EDITOR

Randi Silberman Klett

ASSOCIATE ART DIRECTOR

Erik Vrielink

ASSOCIATE EDITORS

Ariel Bleicher, a.bleicher@ieee.orgRachel Courtland, r.courtland@ieee.orgEliza Strickland, e.strickland@ieee.org

ASSISTANT EDITOR

Willie D. Jones, w.jones@ieee.org

SENIOR COPY EDITOR

Joseph N. Levine, j.levine@ieee.org

COPY EDITOR

Michele Kogon, m.kogon@ieee.org

EDITORIAL RESEARCHER

Alan Gardner, a.gardner@ieee.org

ASSISTANT PRODUCER, SPECTRUM DIGITAL

Celia Gorman, celia.gorman@ieee.org

EXECUTIVE PRODUCER, SPECTRUM RADIO

Sharon Basco

ASSISTANT PRODUCER, SPECTRUM RADIO

Francesco Ferorelli, f.ferorelli@ieee.org

ADMINISTRATIVE ASSISTANTS

Ramona Foster, r.foster@ieee.orgNancy T. Hantman, n.hantman@ieee.org

INTERNS

Davey Alba, Mia Feldman

CONTRIBUTING EDITORS

Evan Ackerman, Mark Anderson, John Blau, Robert N. Charette,

Peter Fairley, Mark Harris, David Kushner, Robert W. Lucky,

Paul McFedries, Prachi Patel, Richard Stevenson, William Sweet,

Lawrence Ulrich, Paul Wallich

DIRECTOR, PERIODICALS PRODUCTION SERVICES

Peter Tuohy

EDITORIAL & WEB PRODUCTION MANAGER

Roy Carubia

SENIOR ELECTRONIC LAYOUT SPECIALIST

Bonnie Nani

SPECTRUM ONLINE

LEAD DEVELOPER

Kenneth Liu

WEB PRODUCTION COORDINATOR

Jacqueline L. Parker

MULTIMEDIA PRODUCTION SPECIALIST

Michael Spector

EDITORIAL ADVISORY BOARD

Susan Hassler, *Chair*; Gerard A. Alphonse, Marc T. Apter, Francine

D. Berman, Jan Brown, Jason Cong*, Raffaello D'Andrea, Kenneth Y.

Goldberg, Susan Hackwood, Bin He, Erik Heijne, Charles H. House,

Chenming Hu*, Christopher J. James, Ruby B. Lee, John P. Lewis,

Tak Ming Mak, Carmen S. Menoni, David A. Mindell, C. Mohan, Fritz

Morgan, Andrew M. Odlyzko, Harry L. Tredennick III, Sergio Verdú,

Jeffrey M. Voas, William Weihl, Kazuo Yano, Larry Zhang*

* Chinese-language edition

EDITORIAL / ADVERTISING CORRESPONDENCE

IEEE Spectrum

733 Third Ave., 16th Floor

New York, NY 10017-3204

EDITORIAL DEPARTMENT

TEL: +1 212 419 7555 FAX: +1 212 419 7570

BUREAU Palo Alto, Calif.; Tekla S. Perry +1 650 328 7570

ADVERTISING DEPARTMENT +1 212 705 8939

RESPONSIBILITY FOR THE SUBSTANCE OF ARTICLES RESTS UPON

THE AUTHORS, NOT IEEE OR ITS MEMBERS. ARTICLES PUBLISHED

DO NOT REPRESENT OFFICIAL POSITIONS OF IEEE. LETTERS TO THE

EDITOR MAY BE EXCERPTED FOR PUBLICATION. THE PUBLISHER

RESERVES THE RIGHT TO REJECT ANY ADVERTISING.

REPRINT PERMISSION / LIBRARIES

Articles may be photocopied for

private use of patrons. A per-copy fee must be paid to the Copyright

Clearance Center, 29 Congress St., Salem, MA 01970. For other

copying or republication, contact Business Manager, IEEE Spectrum.

COPYRIGHTS AND TRADEMARKS *IEEE Spectrum* is a registered

trademark owned by The Institute of Electrical and Electronics

Engineers Inc. Careers, EE's Tools & Toys, EV Watch, Progress,

Reflections, Spectral Lines, and Technically Speaking are

trademarks of IEEE.

CONTRIBUTORS_



Elise Ackerman

An editor at Twilio, a cloud communications company, and a regular contributor to *Forbes.com*, Ackerman last wrote for *IEEE Spectrum* in January, about Google Glass. In this issue, she reports on new workplace BYOD policies [see "The Bring-Your-Own-Device Dilemma," p. 20]. "I remember when my colleagues first started clamoring to bring their iPhones to work," she says. "They had no idea that what they were really proposing was a kind of Faustian bargain—their gadgets for their privacy."



Antonio Liotta

A professor of network engineering at the Eindhoven University of Technology, in the Netherlands, Liotta spends a lot of time picking the brains of biologists and neuroscientists. He takes inspiration from naturally intelligent systems, including animal swarms and the human brain, in his vision for designing future networks, which he describes in the article "The Cognitive Net Is Coming" [p. 24]. His work has instilled in him a love for all networks, he says, "but only if it's not hotel Wi-Fi."



John MacNeill

In his early career, MacNeill progressed from sketching art for toy manuals to a position as an art director for a small electronics magazine in Massachusetts. "But I was better at doing the art stuff than managing the art stuff," he says. Today, the veteran illustrator specializes in using public and classified sources to create 3-D renderings of air, space, and water vehicles, like the *USS Zumwalt* seen in this issue ["Clad in Controversy," p. 30]. "It's like solving a puzzle," MacNeill says.



Jonathan Morris

In 2009, Morris was among those who found single magnetic charges—magnetic monopoles—inside crystals called spin ices. In "The Race for the Pole" [p. 42], he writes about the discovery and what's happened since. Soon to join the faculty at Xavier University, in Cincinnati, Morris was at the time a postdoc at the Helmholtz Centre Berlin for Materials and Energy. The flurry of media attention that ensued included a mention on the U.S. sitcom "The Big Bang Theory." "It was kind of surreal," he says.



Michele Traverso

Traverso, a writer based in Shanghai, says he got the "airplane bug" when he was 4 years old. Now a licensed glider pilot, he has been following the electrification of such aircraft intently, so "Five Electric Planes to Watch" [p. 7] was a perfect fit as his first assignment for *Spectrum*. The designers, engineers, and enthusiasts involved in this effort are "visionaries and dreamers," says Traverso. "Everybody wants [electrification] to work, because it could be so much better than the gasoline engine."



IEEE MEDIA

SENIOR DIRECTOR; PUBLISHER, IEEE SPECTRUM

James A. Vick, j.vick@ieee.org

ASSOCIATE PUBLISHER, SALES & ADVERTISING DIRECTOR

Marion Delaney, m.delaney@ieee.org

RECRUITMENT AND LIST SALES ADVERTISING DIRECTOR

Michael Buryk, m.buryk@ieee.org

BUSINESS MANAGER Robert T. Ross

IEEE MEDIA/SPECTRUM GROUP MARKETING MANAGER

Blanche McGurr, b.mcgurr@ieee.org

INTERACTIVE MARKETING MANAGER

Ruchika Anand, r.anand@ieee.org

LIST SALES & RECRUITMENT SERVICES PRODUCT/

MARKETING MANAGER Ilija Rodriguez, i.rodriguez@ieee.org

REPRINT SALES +1 212 221 9595, EXT. 319

MARKETING & PROMOTION SPECIALIST Faith H. Jeanty, f.jeanty@ieee.org

SENIOR MARKETING ADMINISTRATOR Simone Darby, simone.darby@ieee.org

MARKETING ASSISTANT Quinona Brown, q.brown@ieee.org

RECRUITMENT SALES ADVISOR Liza Reich +1 212 419 7578

ADVERTISING SALES +1 212 705 8939

ADVERTISING PRODUCTION MANAGER Felicia Spagnoli

SENIOR ADVERTISING PRODUCTION COORDINATOR

Nicole Evans Gymah

ADVERTISING PRODUCTION +1 732 562 6334

IEEE STAFF EXECUTIVE, PUBLICATIONS Anthony Durniak

IEEE BOARD OF DIRECTORS

PRESIDENT Peter W. Staecker, president@ieee.org

+1 732 562 3928 FAX: +1 732 465 6444

PRESIDENT-ELECT Roberto de Marca

TREASURER John T. Barr SECRETARY Marko Delimar

PAST PRESIDENT Gordon W. Day

VICE PRESIDENTS

Michael R. Lightner, Educational Activities; Gianluca Setti, Publication Services & Products; Ralph M. Ford, Member & Geographic Activities; Karen Bartleson, President, Standards Association; Robert E. Hebner, Technical Activities; Marc T. Apter, President, IEEE-USA

DIVISION DIRECTORS

Cor L. Claeys (I); Jerry L. Hudgins (II); Douglas N. Zuckerman (III); Jozef Modelski (IV); James W. Moore (V); Bogdan M. Wilamowski (VI); Cheryl "Cheri" A. Warren (VII); Roger U. Fujii (VIII); Jose M. Moura (IX); Stephen Yurkovich (X)

REGION DIRECTORS

Peter Alan Eckstein (1); Parviz Famouri (2); David G. Green (3); Karen S. Pedersen (4); James A. Jefferies (5); Michael R. Andrews (6); Keith B. Brown (7); Martin J. Bastiaans (8); Gustavo A. Giannattasio (9); Toshio Fukuda (10)

DIRECTORS EMERITUS Eric Herz, Theodore W. Hissey

IEEE STAFF

EXECUTIVE DIRECTOR & COO James Prendergast
+1 732 502 5400, james.prendergast@ieee.org

HUMAN RESOURCES Betsy Davis, SPHR
+1 732 465 6434, e.davis@ieee.org

PUBLICATIONS Anthony Durniak
+1 732 562 3998, a.durniak@ieee.org

EDUCATIONAL ACTIVITIES Douglas Gorham
+1 732 562 5483, d.gorham@ieee.org

MEMBER & GEOGRAPHIC ACTIVITIES Cecelia Jankowski
+1 732 562 5504, c.jankowski@ieee.org

STANDARDS ACTIVITIES Konstantinos Karachalios
+1 732 562 3820, constantin@ieee.org

GENERAL COUNSEL & CHIEF COMPLIANCE OFFICER
Eileen Lach, +1 212 705 8990, e.m.lach@ieee.org

CORPORATE STRATEGY Matthew Loeb, CAE
+1 732 562 5320, m.loeb@ieee.org

CHIEF MARKETING OFFICER Patrick D. Mahoney
+1 732 562 5596, p.mahoney@ieee.org

CHIEF INFORMATION OFFICER Alexander J. Pasik, Ph.D.
+1 732 562 6017, a.pasik@ieee.org

CHIEF FINANCIAL OFFICER Thomas R. Siebert
+1 732 562 6843, t.siebert@ieee.org

TECHNICAL ACTIVITIES Mary Ward-Callan
+1 732 562 3850, m.ward-callan@ieee.org

MANAGING DIRECTOR, IEEE-USA Chris Brantley
+1 202 530 8349, c.brantley@ieee.org

IEEE PUBLICATION SERVICES & PRODUCTS BOARD

Gianluca Setti, Chair; John B. Anderson, Robert L. Anderson, John Baillieu, Silvio E. Barbin, Herbert S. Bennett, Don C. Bramlett, Stuart Bottom, Thomas M. Conte, Samir M. El-Ghazaly, Sheila S. Hemami, Lawrence O. Hall, David A. Hodges, Donna L. Hudson, Elizabeth T. Johnston, Hulya Kirkici, Khaled Letalief, Carmen S. Menoni, William W. Moses, Michael Pecht, Vincenzo Piuri, Sorel Reisman, Jon G. Rokne, Curtis A. Siller, Ravi M. Todi, H. Joel Trussell, Leung Tsang, Timothy T. Wong

IEEE OPERATIONS CENTER

445 Hoes Lane, Box 1331, Piscataway, NJ 08854-1331 U.S.A.
Tel: +1 732 981 0060 Fax: +1 732 981 1721

Electric Vehicles Need More Study, Less Emotion

If you've never touched a live wire, try writing about plug-in electric cars

Iur latest such shock came from last month's cover story, *Unclean at Any Speed*, as evidenced in the astounding volume of comments it elicited. As author Ozzie Zehner notes, "the seemingly simple question 'Are electric cars indeed green?' quickly gets complicated." • The question is twofold: How much carbon does a car emit in operation, and how much carbon is emitted in making the car—and its materials?

The first problem involves the well-to-wheel carbon footprint—the carbon emitted by obtaining fossil fuels, converting them into electricity, and using that electricity to produce motive force. Two U.S. reports have compared electric and gasoline cars on that basis, the more recent one coming from the Union of Concerned Scientists.

The study finds that in all U.S. states, plug-in electric vehicles are at least as fuel-efficient as the best gasoline cars. In many states, plug-in EVs are as good as the best gasoline-electric hybrids. And in some, they're better than the best hybrids. Gasoline hybrids beat pure electrics only in parts of the country that generate most of their power from coal.

An earlier, more comprehensive study was issued jointly by the Electric Power Research Institute (EPRI) and the Natural Resources Defense Council (NRDC) in 2007. The conclusion: Electrifying transportation would reduce greenhouse gases, and as the grid gets cleaner, the benefits increase. So score one for the electric car.

The next place to look is at materials and manufacturing. That's where last month's cover story comes in. It argues that making lithium-ion battery packs, electric motors, ultralight materials, and power electronics releases more carbon into the atmosphere than the vehicle saves in operation. But just how important are materials and manufacturing? According to a 2000 study by the Energy Laboratory at MIT, extraction of the raw materials to build the vehicle makes up just 4 percent of a vehicle's lifetime carbon footprint, and building it adds another 2 percent. Zehner relies instead on a 2010 study by the National Research Council of the National Academies of Sciences, whose numbers differ greatly from MIT's. The NAS concluded that the lifetime health

and environmental damages from electric cars exceed those from gasoline cars. If nothing else, this highlights the urgent need for more studies on the environmental impacts of materials and manufacturing.

But the important thing to remember is that expensive new things usually get cheaper as the volume of production rises. So the carbon burden of making electric cars will likely fall in the future—even as it rises for conventional cars. Carmakers are working hard to use less of the costly metals, and even eliminate rare earth metals. Already Tesla Motors uses none of those metals in its motor or battery. In any case, manufacturers of electric cars are not alone in their use of lighter materials—conventional cars will increasingly have them too. One example: A Tesla Model S is mostly made of aluminum, but so are the current Audi A8, Jaguar XJ, and Range Rover.

The media coverage, gnashing of teeth, shrieking about conspiracies, and activism by advocates might lead readers to believe that electric cars are either the solution to many of our ills or a boondoggle on wheels. Zehner notes this polarization; I would add that the extremes often seem to emanate from politically partisan sources.

But the more optimistic view would be that such vituperation merely indicates that we are in the early stages of the "change curve," based on work by the psychiatrist Elisabeth Kübler-Ross. It starts with denial, progresses to anger, moves into exploring, and finally reaches acceptance. Already, with more than 100 000 electric cars on U.S. roads, we've moved from denial to anger. Now it's time to go into the labs and explore. —JOHN VOELCKER

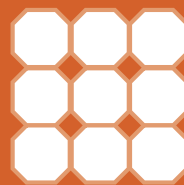
John Voelcker is editor of *Green Car Reports* and a senior editor at High Gear Media. He covers auto technologies for both consumer and industry outlets.

Because of the great interest in this topic, we've created a special section online. There you'll find additional commentary on EV research from Mark Duval, director of electric transportation and energy storage at the EPRI, the IEEE Transportation Electrification Initiative Steering Committee, and others (<http://spectrum.ieee.org/evcars0813>).

Note that IEEE, as a whole, does not typically take official positions on technology topics, and that articles appearing in *IEEE Spectrum* represent the thoughts of our authors, not IEEE. *Spectrum's* editorial mission is to act as a forum for the discussion of new and emerging technologies.



NEWS

**SOLAR IMPULSE CARRIES
11 000 SOLAR CELLS****ARE WE THERE YET?** The Solar Impulse photovoltaic-powered plane took a break in this St. Louis hangar about halfway through its cross-country U.S. tour.

FIVE ELECTRIC PLANES TO WATCH

**Solar and other all-electric
planes are getting more practical**



SOLAR IMPULSE, A 1-SEATER SOLAR-POWERED PLANE, completed a 5695-kilometer bunny hop across the United States on 6 July. The feat deservedly got a lot of attention, but it's just the tip of the iceberg: Electric flying has been a reality for quite some time, and it's never been more practical.

Aviation is a slow-moving industry, but the daring designers of electric aircraft have made a lot of progress recently. During the past two years, as a number of key technologies—batteries, controllers, motors, and materials—have neared maturity and become easier to source at more affordable prices, there has been a minor boom of sorts in the offering of electric drives for small planes.

Because batteries still haven't made the energy-density quantum leap that we all hoped for—gasoline's energy density is still about 13 times that of the best lithium-ion batteries—most electric planes are self-launching gliders or motorized gliders. These have less stringent requirements in terms of range than standard aircraft, and their highly aerodynamically efficient airframes require less power to keep them airborne in all phases of flight. What follows is a sampling of the most innovative efforts. There are more at <http://spectrum.ieee.org/10electricplanes>.

SOLAR IMPULSE

The recent Solar Impulse flight has been both dismissed as a pointless exercise in renewable energy technology and praised as a daring Victorian-style adventure with exciting technology implications. Whichever side you take, the facts are that the people behind Solar Impulse pulled an impressive team together, managed to build a unique solar-powered airplane, and flew it first from Switzerland to Morocco and then across the United States.

Solar Impulse (tail number HB-SIA) is the first of two aircraft planned by the Switzerland-based group. As a test bed for numerous more or less innovative technologies and construction techniques, the HB-SIA is a collection of impressive numbers: It has the wingspan of a major airliner (about 63 meters), but its honeycomb-shaped, carbon-fiber-reinforced structure keeps the weight down to that of a small car (1600 kilograms); a fourth of that weight comes from the batteries. That low weight and its sleek design mean it can get by on four 7.35-kilowatt motors. It's not built to be practical, and it isn't: It flies at an average speed of 70 kilometers per hour, and it fears rain and winds that any student pilot in almost any plane could handle without too much worry.

Despite its pedigree, the fate of HB-SIA is sealed: The cross-U.S. trek will be its last. A new plane is being built in Switzerland—the HB-SIB. It will incorporate all the lessons learned over these past few years of intensive flight testing and publicity tours, in preparation for the real goal of Solar Impulse founders Bertrand Piccard and André Borschberg: a round-the-world flight in the spring of 2015. Godspeed, gentlemen.

ANTARES 23E

With its 23-meter wing—designed by the Dutch aerodynamics guru Loek Boermans and built by Germany's Lange Aviation—the Antares 23E is the most advanced electric airplane money can buy. The new wing profiles, along with the newly optimized wing/fuselage interface, give the Antares 23E a glide ratio of 60—meaning that from a height of 1000 meters the Antares 23E can glide for 60 km. The plane is an impressive machine, built with safety in mind: It's one of the few production gliders equipped with a crash-proof cockpit, and the design and manufacturing process of each component has been thoroughly assessed by German aviation authorities. In fact, it's the world first electric aircraft to receive certification by any aviation authority.

The aircraft's motor and battery pack haven't changed much from those of an earlier 18-meter-wingspan version. It has the same 42-kW brushless motor and the same SAFT VL41M lithium-ion batteries used on most European satellites and the F35 Joint Strike Fighter. However, by using a better battery-charging process, its engineers managed to squeeze 500 extra meters of climb out of it, bringing the maximum altitude to 3500 meters on a single charge. This albeit excellent piece of German engineering will set you back €205 000 (about US \$260 000), onboard charger included.

TAURUS ELECTRO G2

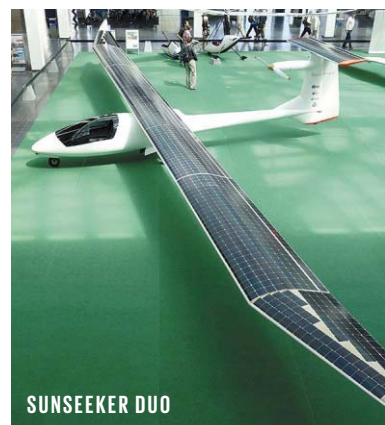
Pipstrel is a smallish Slovenian airplane maker that has managed to climb to global relevance thanks to the relentless efforts of its founder, Ivo Boscarol. Its line of efficient, prizewinning light aircraft was completed in 2007 with the launch of the Taurus, a side-by-side-seat self-launching light glider. The Taurus G4 prototype, with its daring twin-fuselage design, won the prestigious 2011 Google Green Flight Challenge. Its commercial electric incarnation, the G2, launched the same year, features a 40-kW engine and a Li-ion battery pack that lets it climb 2000 meters on a single charge. What's more, it



SOLAR IMPULSE



FRONT ELECTRIC SUSTAINER



SUNSEEKER DUO



ANTARES 23E



TAURUS ELECTRO G2

comes with a solar-panel-covered trailer that can recharge the batteries in 5 hours.

Pipistrel should be credited for stating clearly that building an electric version of the Taurus was not an effort to go green but rather a technological improvement: The electric plane climbs faster and has a shorter takeoff run than does its petrol-powered version. As if this wasn't enough, the enterprising Slovenian company has designed a highly streamlined new 4-seater called Panthera. The plane made its first flight this spring on a combustion engine, but Pipistrel plans to offer this sleek airplane in both hybrid and pure electric versions, with a 145-kW engine.

SUNSEEKER DUO

The Sunseeker family of solar gliders dates back almost 30 years. Developed and flown by the talented Californian engineer Eric Raymond, the Sunseeker I crossed the United States in the summer of 1990. Raymond, a protégé of solar flight legend Paul MacCready, took 21 flights and 121 hours in the air to make that historic crossing. After an upgrade of the Sunseeker's electronics in 2009, Raymond decided to develop a whole new aircraft, this time with two seats instead of one. With solar panels carefully placed on the surface of the wings and on the tail, the Sunseeker Duo can cruise under the power generated by its own cells. It has some unique, elegant features, such as a sliding canopy that provides the best in-flight views in any class, swiveling tablet support for the next-generation avionics and maps running on the iPad, and—wait for it—seats that recline completely. So go ahead, stretch out and snooze if you like. Just make sure you have an alert copilot.

FRONT ELECTRIC SUSTAINER PROPULSION

Another small but clever Slovenian company (we sense a trend), LZ Design had what now seems like a perfectly obvious idea. Rather than design complicated mechanisms for a retractable mast with an engine and a foldable propeller that sits in a bay behind the pilot, the company decided to put an electric engine in the nose of the glider and the batteries in the compartment that's usually reserved for a small two-stroke engine, as is found in some of the world's best self-launching gliders.

The brushless engine drives an ultralight carbon propeller (240 grams), whose blades fold elegantly along the curvature of the fuselage and fit snugly into thinly molded recesses when not in use. The drag of the hinge is as negligible as that of an insect wiper. (Yes, glider pilots are so obsessed with aerodynamic efficiency that they keep their wings clear of squashed bugs to reduce drag.)

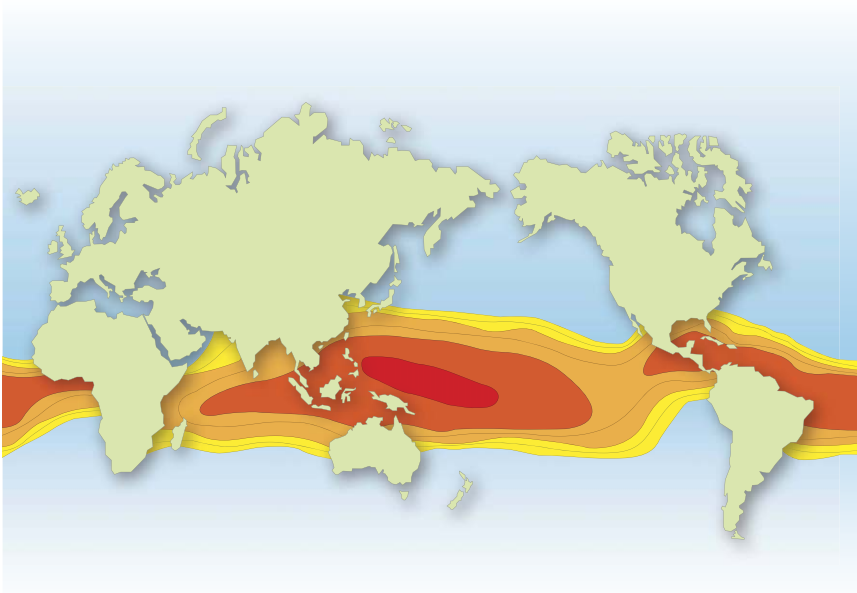
A front electric sustainer propulsion system is installed on a few single-seaters, such as the light Italian glider Silent 2 Electro and the Lithuanian LAK 17. The company is searching for alternative energy storage options.

—BY MICHELE TRAVERSO

NEWS

LOCKHEED MARTIN PIONEERS OCEAN ENERGY IN CHINA

A 10-megawatt ocean thermal energy conversion plant is under way



Just a few years ago, Lockheed Martin was working to build a pilot plant to demonstrate a renewable energy

technology called ocean thermal energy conversion (OTEC) near the Hawaiian island of Oahu. The company wanted to get funding from the U.S. Navy for the pioneering project and to cable the electricity it produced straight to the naval base at Pearl Harbor.

Now Lockheed is designing that 10-megawatt pilot plant—but not in American waters. Instead, the facility will be off the coast of southern China, and Lockheed's customer is a private Chinese company that develops resorts and luxury housing.

Over the years Lockheed has approached various potential partners, says Rob Varley, the company's OTEC project manager. Building an offshore energy station at commercial scale is an expensive proposition, particularly when it's the first time the technology is being tried out. Lockheed won't release the cost of the project, but outside experts estimate that a 10-MW facility would cost roughly US \$300 million to \$500 million. However, experts say that a full-scale 100-MW plant would be more competitive at just \$1.2 billion.

"The biggest challenge has been to get the gold and start the project," says Varley, but in terms of engineering, he says, "I don't see any showstoppers at this point."

HOT WATER: Areas with the greatest temperature differential between warm surface waters and cold deep waters [red], are best suited for ocean thermal plants.

That's not surprising, since the company has been working on OTEC since the 1970s, and the technology hasn't changed drastically since then. OTEC systems make use of the temperature differential in tropical areas between warm surface water and cold deep water. In most systems, ammonia, which has a very low boiling point, passes through a heat exchanger containing the warm water. The ammonia is vaporized and used to turn a turbine, and then it's cycled past the cold water to recondense. This is a renewable energy technology with the rare capacity to supply base-load power, as water temperatures are fairly stable.

The ammonia passes through a closed loop, while the water comes and goes through massive pipes. The project in China may pump cold water up from a depth of about 1000 meters, using a pipe that's 4 meters across. Varley says that some of the infrastructure can be borrowed from the offshore drilling industry: "We showed them our requirements for the platform, and they yawned and said, 'Is that all you got?'" he says. "But then we showed them the pipe." Attaching the massive pipe to a relatively small floating platform creates unusual stresses, Varley says. Lockheed also had to find materials for the pipes and the heat exchangers that could withstand the harsh marine environment.

Lockheed's client is Reignwood Group, a Chinese company whose diverse portfolio includes resort and housing developments. According to a company press release, Reign-

wood Group wants the 10-MW plant to supply all the power for a large-scale environmentally sound resort community that the company will build in southern China. A Reignwood spokesperson did not respond to requests for more details by press time. A Lockheed spokesperson says the companies are currently working on site selection and that they'll start designing a facility this year to suit the specific conditions at that site.

The China project isn't the only OTEC project going ahead. Baltimore-based OTEC International is negotiating the terms of a 1-MW demonstration plant in Hawaii, and the company is planning much bigger facilities in Hawaii and the Caribbean. Both OTEC International and Lockheed Martin see their current plans as steps toward a much more ambitious goal: utility-scale OTEC plants. "Going from a PowerPoint to a 100-MW would be too big a leap," says Lockheed's Varley.

OTEC advocates have been trying to build megawatt-scale facilities for decades, but several ambitious projects have failed to materialize. So why should it be different this decade? Eileen O'Rourke, president of OTEC International, says there's a convergence of favorable conditions. "Island jurisdictions like Hawaii have very high energy prices and limited alternatives for base-load power, and OTEC fits with their desire to be energy independent and green," she says. Add in mature technology from the offshore oil industry, she says, and "we just think the time is right for OTEC." —ELIZA STRICKLAND

Some experts say that a 10-megawatt ocean thermal energy conversion plant would cost US \$300 million. A much larger, 100-MW facility would require \$1.2 billion, they estimate.

LONG-DISTANCE QUANTUM CRYPTOGRAPHY

A hybrid system could extend quantum cryptography's reach



Using the quirky laws of

quantum physics to encrypt data, in theory, assures perfect security. But today's quantum cryptography can secure point-to-point connections only about 100 kilometers apart, greatly limiting its appeal.

Battelle Memorial Institute, an R&D laboratory based in Columbus, Ohio, is now building a "quasi-quantum" network that will break through that limit. It combines quantum and classical encryption to make a network stretching hundreds of kilometers with security that's a step toward the quantum ideal.

"In a few years, our networks aren't going to be very secure," says Don Hayford, senior research leader in Battelle's national security global business. Cryptography relies on issuing a secret key to unlock the contents of an encrypted message. One of the long-standing worries is that sufficiently powerful computers, or eventually quantum computers, could decipher the keys. "We looked at this and said, 'Somebody needs to step up and do it,'" Hayford says.

By the end of next year, Battelle plans to have a ring-shaped network connecting four of its locations around Columbus—some of which transmit sensitive defense contract information—that will be protected using quantum key distribution, or QKD. If



that smaller network is successful, Battelle then plans to connect to its offices in the Washington, D.C., area—a distance of more than 600 km—and potentially offer QKD security services to customers in government or finance over that network.

Quantum cryptography uses physics, specifically the quantum properties of light particles, to secure communications. It starts with a laser that generates photons and transmits them through a fiber-optic cable. The polarization of photons—whether they're oscillating horizontally or vertically, for example—can be detected by a receiver and read as bits, which are used to generate the same “one-time pad” encryption key at both ends of the fiber. (A one-time pad is an encryption key that consists of a long set of random numbers, and so the message it hides also appears to be a random set of numbers.) Messages can then be sent securely between the sender and receiver by any means—even carrier pigeon—so long as they are encrypted using the key. If someone tries to intercept the key by measuring the state of the photons or by reproducing them, the system will

be able to detect the intrusion and the keys will be thrown out.

Over long distances, though, light signals fade, and keys can't be distributed securely. Ideally, “quantum repeaters” would store and retransmit photons, but such devices are still years away, say experts. Battelle's approach is essentially to daisy-chain a series of QKD nodes and use classical encryption to bridge the gaps. Locations less than 100 km away will be connected by fiber-optic links and the data secured by a QKD system from Geneva-based ID Quantique. For two more-distant nodes (call them A and C) to communicate, there must be a “trusted node” between them (call it B). Nodes A and B can share a key by quantum means. Nodes B and C can also share a separate key by quantum means. So for A and C to communicate securely, A's key must be sent to C under the encryption that B and C share. You might think the quantum-to-classical stopover in the trusted node might be a weak point, but even inside that node, keys are protected using one-time pad encryption, says Grégoire Ribordy, the CEO and cofounder of ID Quantique. The trusted node will also

be located at a secure site and have other measures to prevent tampering.

These nodes, which are still under development, will be designed to integrate with corporate security systems, distributing keys for virtual private networks or database security within a building. “The idea is to set up a network which would be dedicated to cryptography-key management,” says Ribordy. ID Quantique's gear will do the quantum key exchange, while Battelle will build the trusted nodes.

Researchers also hope to treat satellites in space as trusted nodes and to send photons through the air, rather than over optical-fiber links. In the nearer term, though, Battelle's land-based QKD network may be the most viable approach to introducing quantum encryption into today's networks. Yet it still faces significant challenges. For starters, the cost of point-to-point QKD is about 25 to 50 percent more than for classical encryption,

says Ribordy, and connecting locations hundreds of kilometers apart would require multiple systems. That means Battelle will need to find a customer with an application that warrants the added expense. Verizon Communications, which offers network security services, tested QKD from 2005 to 2006, but it determined there wasn't a viable business case because of distance limitations and the limited market for the technology.

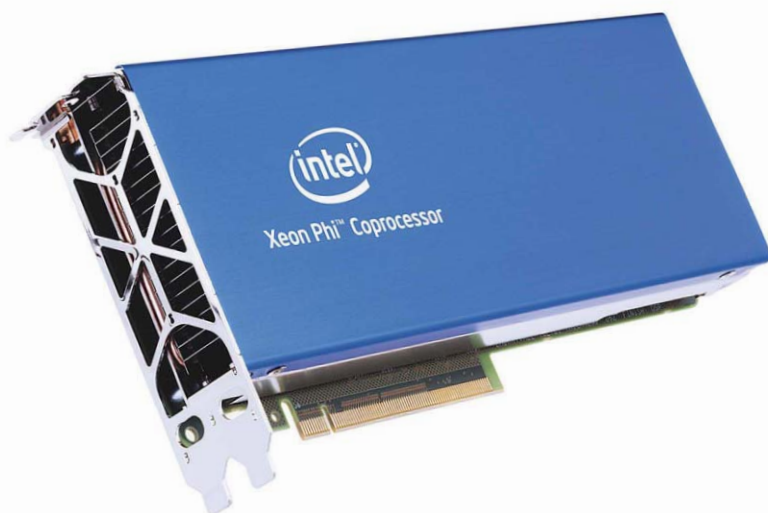
Also, QKD hardware can't easily plug into the existing telecom hardware, says Duncan Earl, chief technology officer of GridCom Technologies, which plans to use QKD for electricity grid control networks. Established networks have routers and switches that would ruin the key distribution's delicate physics.

On a technical level, though, the work really only requires good engineering, not scientific breakthroughs, says Hayford. And the hybrid approach can accommodate future advances in quantum cryptography, such as quantum repeaters. Given the growing concerns over cybersecurity, it's better to test the worth of quantum encryption sooner rather than later, he says. —MARTIN LAMONICA

MICHAEL BODMANN/GETTY IMAGES

NEWS

NEWS



INTEL STRIKES BACK

Its new massively parallel chips take the lead in the latest Top500 supercomputer rankings



➤ **When computer scientist Jack Dongarra** traveled to China in May to see the Tianhe-2 supercomputer, he was ready to be impressed. And the machine, built by the National University of Defense Technology, in Changsha, didn't disappoint. By early June, Tianhe-2 had demonstrated a peak speed of 33.86 petaflops: some 34 million billion floating point operations per second, far more than what it needed to place first on June's Top500, a biannual list of the fastest supercomputers, compiled and adjudicated by Dongarra and his colleagues.

But Tianhe-2 isn't just fast; it could also be a bellwether for the field of scientific supercomputing. The Intel-based machine is a hybrid that uses 32 000 multicore central processing units (CPUs) and 48 000 "coprocessors"—separate "accelerator" chips that incorporate dozens of cores that are specially designed to churn through the floating-point arithmetic operations needed for simulations of climate and the early universe, among many others.

Advanced Micro Devices (AMD) and Nvidia Corp., which make general-purpose graphics-processing units (GPGPUs) that are well-suited for this sort of application, have previously dominated the accelerator market. But Intel is making a strong bid for the

space, with a different kind of accelerator it calls the Xeon Phi coprocessor.

Seven systems on the Top500 list were already using the Intel chip by the time it made its official debut in November 2012. Among top supercomputers, Intel's accelerator share is still small; machines using Nvidia's GPUs outnumber those with Intel's coprocessors 31 to 11. But Tianhe-2 has had a big impact on the landscape. If you tally up all the petaflops on the list, "the aggregated performance delivered by Intel Xeon Phi coprocessors is now bigger than the performance delivered by GPU accelerators," says Intel spokesperson Radoslaw Walczyk. "It is a big win," says Sergis Mushell, an analyst at the technology research firm Gartner.

Intel's chips contain up to 61 cores and are built using the company's 22-nanometer manufacturing process, which is a generation ahead of the competition. The company says its coprocessors have a few advantages over GPGPUs: They can operate independently of CPUs and they don't require special code to program. But it will take some time to see how Intel's venture will fare.

The market for accelerators is still small. As of June, just 53 of the 500 computers on the list use them, down from 62 in November. But

this small decline may be temporary. "They deliver much more bang for your kilowatt," says Michael Shuey, a systems architect in high performance computing at Purdue University. And even though just a tenth of the machines on the list incorporate accelerators, Dongarra says, the machines account for a third of the aggregated performance. "It's a small number that has a large impact," he says.

One stumbling block to more widespread adoption has been programming complexity, says William Gropp at the University of Illinois at Urbana-Champaign. The university hosts the U.S. National Science Foundation's most powerful machine, Blue Waters. The CPU- and GPU-based system sustains speeds of more than 1 petaflops.

As with any machine that pairs CPUs and accelerators, Blue Waters' programmers must find effective ways to transport data between the chips, a process that can consume a lot of power. "Anything you do has to be a big enough lump of work in order to amortize moving the data over," Gropp says. This makes programming more difficult: "You have to figure out how to aggregate [work] into a lump whose size is not based on the characteristics of your problem but the characteristics of your hardware."

"We're making significant strides in dealing with the added complexity," Gropp adds. "In the short term, you'll see more and more of these systems with accelerators. [But] I don't think that they'll completely populate the list." CPU-only supercomputers, such as IBM's Blue Gene machines, he says, may remain competitive for a while yet.

What could ultimately take over supercomputing are computers that have both CPUs and specialized cores on the same chip. AMD, Intel, and Nvidia make such chips for mobile devices and personal computers. Although this sort of "heterogeneous integration" has yet to emerge in high-performance systems, Intel says a version of its next coprocessor, Knight's Landing, could be installed directly into supercomputer motherboard sockets. For some computer scientists, integration is just a matter of time. "We're going to see things get closer together," says Dongarra. "It's a natural trend."

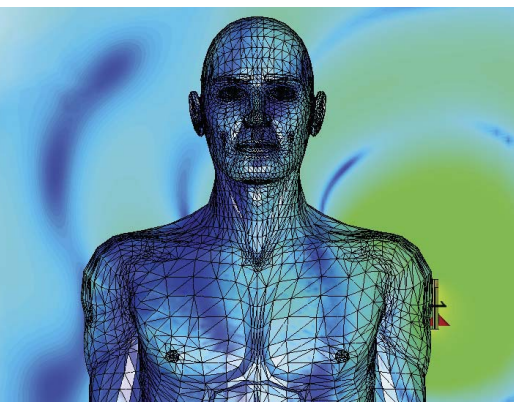
—RACHEL COURTLAND

INTEL



Make the Connection

Find the simple way through complex
EM systems with CST STUDIO SUITE



Components don't exist in electromagnetic isolation. They influence their neighbors' performance. They are affected by the enclosure or structure around them. They are susceptible to outside influences. With System Assembly and Modeling, CST STUDIO SUITE helps optimize component and system performance.

Involved in antenna development? You can read about how CST technology is used to simulate antenna performance at www.cst.com/antenna.

If you're more interested in filters, couplers, planar and multilayer structures, we've a wide variety of worked application examples live on our website at www.cst.com/apps.

Get the big picture of what's really going on. Ensure your product and components perform in the toughest of environments.

**Choose CST STUDIO SUITE –
Complete Technology for 3D EM.**







MOUNTAIN MONITOR

Some of the scores of skiers, hikers, and climbers who die each year attempting to traverse the Alps might have survived if rescuers had reached them in time. But sending in the cavalry without up-to-the-minute information could endanger the would-be saviors. Drone-based reconnaissance now seems practical, after an autonomous MD4-1000 quadcopter made by German drone manufacturer Microdrones braved the adverse conditions along the Alps' Saint-Gotthard Massif mountain range in June. The 25-minute, 12-kilometer trip from Hospental to Airolo, Switzerland, was far from a walk in the park. The 5.6-kilogram microdrone's path was preprogrammed with 18 GPS waypoints selected to help it avoid obstacles such as power and telephone lines, but its minders still had to make a few adjustments to ensure that severe weather didn't cause it to meet the same fate as Hannibal's elephants did.

THE BIG PICTURE

NEWS

We're making waves

With the world's most advanced EM and RF simulation software and measurement systems, it's no wonder we're getting noticed.

From concept to design to fielded solutions, Delcross has the tools to help you ride the wave of success. Our Savant software rapidly solves installed antenna problems that other electromagnetic solvers cannot address. With GPU acceleration in Savant, the power of a cluster can be realized for the price of a video card. Using our EMIT software and AMS measurement system, you can extend your analysis beyond the antenna ports to tackle your toughest cosite interference and link margin analysis problems. We provide unique solutions that get you ahead.

www.delcross.com



DEL CROSS
TECHNOLOGIES

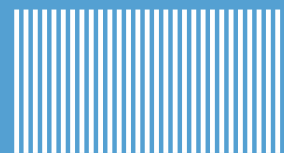
For more information, please contact your local CST AG representative (www.cst.com).

Savant

EMIT

AMS

RESOURCES



30 LINES: THE RESOLUTION OF THE FIRST
COMMERCIALY PRODUCED TELEVISION,
THE 1930 BAIRD TELEVISOR.



RESOURCES_TOOLS

MOBILE TV
WATCHING
BROADCAST
TV ON YOUR
iPHONE

Even as more and more viewers shift to getting television programs delivered via fixed and mobile Internet connections, broadcast television still has its place. Over-the-air transmissions are more resilient than a video stream. There's no network congestion or overwhelmed servers during a live event, for example, and when disasters rob cell towers or home routers of power, television stations can still reach affected areas. So it's perhaps a little bit surprising that the options for watching broadcast TV on mobile devices have been thin on the ground, which makes gadgets like Escort's MobileTV for iOS devices very welcome. • The dearth of such receivers is not for lack of effort: Televisions small enough to carry in your pocket have been around since the 1970s, but the early devices used miniaturized cathode ray tubes, which were pricey and power hungry. The proliferation of low-cost touch screens, powerful mobile processors, and digital TV broadcasts should have made things easy, but the original HDTV broadcast standards posed a problem. • Developed in the early to mid-'90s, these standards assumed that television sets would remain in one place when in use. Consequently, phenomena such as Doppler shifting of the broadcast signal can disrupt reception when the receiver moves around. In 2009, a mobile-friendly version of the U.S. digital TV standard was released. However, deployment has been slow, because it requires equipment upgrades that many broadcasters were reluctant to make so soon after the major expense of switching from analog to digital transmissions. The upshot has been that current mobile TV solutions, such as those offered by Aereo, receive the broadcast signal with a fixed antenna and then relay that signal to users via cellular or Wi-Fi networks. The rollout of mobile-friendly digital TV is still a work in progress, but so far 42 cities now have at least one upgraded station. ▶

PHOTOGRAPH BY TAKA

SPECTRUM.IEEE.ORG | INTERNATIONAL | AUG 2013 | 17

RESOURCES_HANDS ON

Escort's US \$100 MobileTV is a matchbox-size unit that takes advantage of the new system. It plugs directly into iOS devices that use the original 30-pin connector socket; owners of iPhones and iPads with the new Lightning socket will have to use an adapter. It comes with a small telescoping antenna and a USB charging cable (the MobileTV has its own battery). Setup is a matter of downloading a free app and scanning for upgraded television stations; in the Boston area I'm able to pick up three. The app incorporates basic DVR functions, for example, allowing viewers to rewind live TV.

I tested the MobileTV during a train ride on the Acela Express service between Boston and New York City. This train can reach speeds of up to 240 kilometers per hour. Not surprisingly, reception suffered when the train went through culverts or tunnels, but it appeared unaffected by the movement of the train itself, with good picture and sound quality. The signals faded completely a little beyond 23 km from Boston's city center. As there are no upgraded stations between Boston and New York, it wasn't until I drew close to the Big Apple that I could begin to watch television again, this time about 50 km away from the city center.

For me, this device feels like the perfect addition to the "urban survival kit" I carry with me in my messenger bag—a multitool, an LED flashlight, and a battery-based USB charger—all of which I have found myself in need of on occasion while making my way around New York and Boston. The ability to watch broadcasts while out and about during times such as the aftermath of an explosion or power failure, instead of trying to glean information byte by byte via a smartphone browser, is an appealing one indeed.

—STEPHEN CASS

FORGET-ME-NOT BEACONS

DIMINUTIVE TRANSMITTERS HELP PREVENT YOU FROM LEAVING SOMETHING BEHIND



A

ALL MY LIFE, I'VE MISPLACED THINGS. I'VE STOPPED COUNTING

the number of umbrellas and jackets I've lost. So it's comforting that technology has finally caught up with this common human foible.

Owners of recent-generation iPhones can now purchase a variety of low-power radio beacons that they can attach to their clothing, bags, wallets, or whatever they fear leaving behind. When the item gets out of range, their phones will alert them. But there are places where cellphone use is forbidden, such as schools and locker rooms. And the beacons are a little problematic for keeping track of multiple items, because each one can cost US \$50 or more. Then there are people like me, who don't have an iPhone and don't plan to have one anytime soon. So I decided to build a phoneless loss-prevention system. • Not wanting to start completely from scratch, I decided to see if I could press the Nike + iPod system into service. This uses a \$19 wireless transmitter, designed to be inserted in certain Nike shoes to keep track of your running. • The Nike + iPod system has been around since 2006, and as a consequence talented hackers have had lots of opportunities to learn how it works. A sensor registers your footsteps and transmits that information to a

receiver dongle that you plug into an iPod. You can buy the dongle bundled with a sensor for \$30. (If you have an iPhone 4S or 5, or a fifth-generation iPod Touch, you can dispense with the dongle entirely and receive reports directly from the transmitter in your shoe.)

The Nike + iPod transmitter, it seemed to me, would be the perfect basis for a homebrew radio beacon. But I soon discovered a complication: It is designed to stop transmitting 10 seconds after it ceases detecting movement. Another issue is that the transmitter, like many Apple products, doesn't have a replaceable battery. Once that goes, the thing is toast.

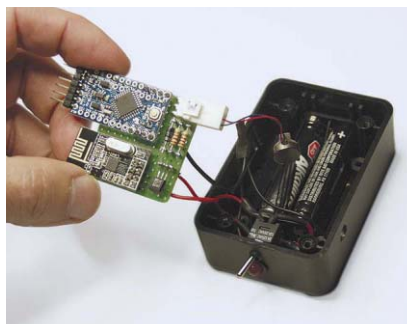
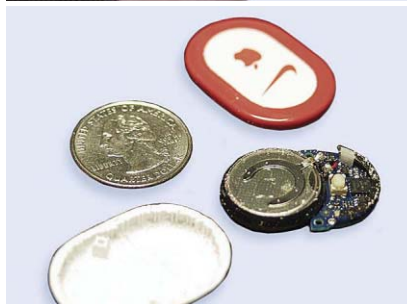
But just maybe there'd be ways around those issues, I thought, as I sawed open the plastic case of the sensor to expose the circuit board inside. It wasn't obvious from first inspection, though, how to keep the board's one-second-interval transmissions repeating indefinitely.

A search online revealed a teardown by *EE Times* that identified one of the ICs on the board as an analog voltage comparator—a Microchip MCP6541. It struck me that this component was surely being used as part of the motion-detection circuitry. Probing that IC's output pin confirmed that the voltage went from about 3 volts (the supply voltage) to zilch whenever the unit was shaken. So I severed the output pin on the comparator IC and connected the pad to which it had been attached to ground. This was easy enough to do despite the tiny size of the IC, as it required only a solder bridge to the next pad over.

Alas, this bit of hacking didn't accomplish my goal. In desperation, I started experimenting with how the solder bridge I'd just made affected the transmitter's one button. Normally this button is used to prevent the circuit from entering its transmit mode—say, if you're carrying it on a train or plane and don't want the ongoing motion to run down the battery. To my delight, I discovered that depressing the button now kept the unit transmitting. So it was straightforward to add two jumpers to permanently short the contacts on this button.

Now I had a beacon that would continue to transmit once a second for as long as the battery held out. According to the Nike + iPod instructions, that should be something like

JUST DOING IT



Starting with an off-the-shelf Nike + iPod transmitter [top], I removed and modified the circuit board [second from top]. I constructed a receiver from an Arduino Pro Mini microcontroller and a Nordic radio chip [second from bottom]. Then I repackaged the transmitter to make it easy to replace its battery and switch it on and off [bottom].

1000 hours, which works out to a little over a month. So I'd be going through a lot of batteries if I left it on all the time. And changing batteries by opening the case and soldering wires to coin cells promised to get old quickly. This got me looking for a better solution than just packaging things back up in the original plastic shell.

The perfect answer appeared in Adafruit Industries' online catalog: a \$2.50 plastic enclosure that holds two coin cells and also houses a switch. I easily modified it to hold just one coin cell and the transmitter circuit board in a neat little package, making for a low-power radio beacon that I could just switch on and off.

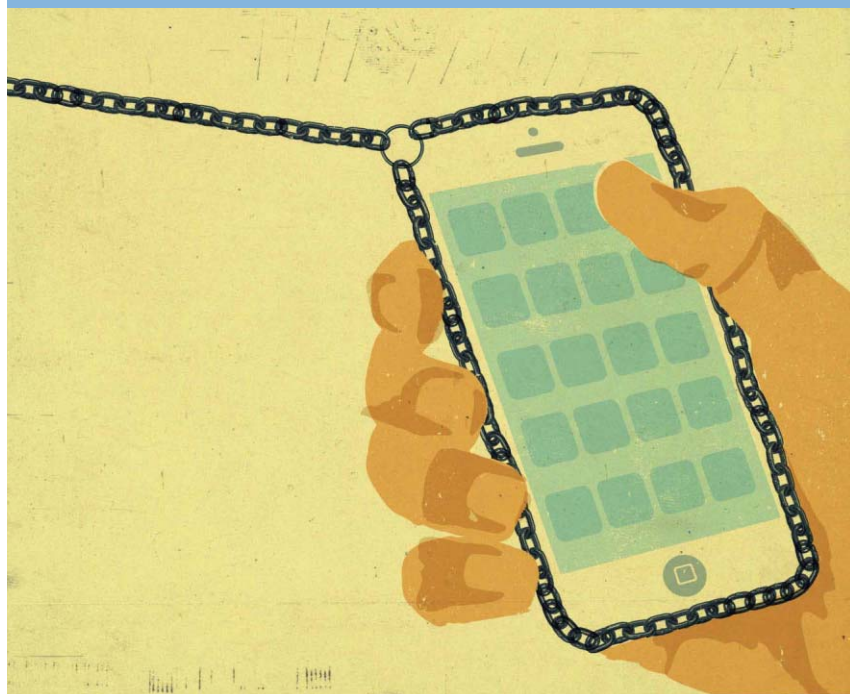
I also sent away to Sparkfun Electronics in Boulder, Colo., for a \$25 adapter board designed to allow you to monitor the output of the Nike + iPod dongle. For a portable receiver, I initially combined this adapter and a dongle with an Arduino Pro Mini microcontroller (\$10), also from Sparkfun. But then I discovered that a programmer named Dmitry Grinberg had done some very clever reverse engineering of this system, which allowed me to monitor the beacon's transmissions in a much cheaper way.

To receive the signals without the dongle, I connected the Arduino Pro Mini to an inexpensive breakout board hosting the Nordic nRF24L01 radio chip (\$6 from Sainsmart.com). After I'd followed the general instructions accompanying the Arduino RF24 software library, I found comments posted on the library developer's website that revealed how to hack the library to receive signals from a Nike + iPod transmitter. So in no time I had a beacon talking to a pocket-size receiver for about \$45 in parts. This was less than one of the commercial beacons alone would cost—and it saved me the expense of an iPhone or fifth-generation iPod Touch. —DAVID SCHNEIDER

RESOURCES_AT WORK

THE BRING-YOUR-OWN-DEVICE DILEMMA

EMPLOYEES AND BUSINESSES SEEK TO BALANCE PRIVACY AND SECURITY



didn't want to be seen as Big Brother. Intel doesn't track personal communications, but employees must sign an agreement that includes a code of conduct and software licensing guidelines, and their devices can be remotely wiped. User training covers unacceptable usage and behavior, such as downloading pirated videos or loaning a device with corporate data to a family member.

Analysts and technology vendors typically make the case that the money saved by businesses in having their employees purchase devices is enough to pay for recommended concomitant investments in security and management infrastructure. But Intel found that wasn't necessarily the case. "When we first started looking at this, we realized we weren't going to save a lot of money," says David Buchholz, Intel's director of consumerization. Three years later, Buchholz says Intel has realized a "soft ROI [return on investment]," measured primarily by employee claims that they save 57 minutes a day by using their own devices.

As IT departments get closer to determining the true cost of BYOD programs, some employees are taking a hard look at the policies and questioning whether the trade-off in personal control is worth it. When the state of Michigan recently prepared to start a pilot BYOD program, the No. 1 concern expressed by users was losing cherished family photos stored on a personal device to a remote wipe. "It's not 'Your Own Device' if you're letting someone else control it," a Slashdot commenter wrote in a discussion of BYOD policies in May. "If I've paid for hardware with my own money, it's mine... No one else gets admin, root, remote-wipe, find my iPhone, or whatever privileges but me."

Forrester's Johnson says he doesn't expect concerns about cost or control to derail BYOD. He notes that new approaches, such as keeping corporate data on the device in a separate software container (which allows the user's and the business's programs to run simultaneously without accessing each other's data), are both more secure and less intrusive. "BYOD is inevitable," he says. All that remains to be worked out are the policy details. —ELISE ACKERMAN

The smartphone revolution opened the floodgates to the BYOD (bring your own device) trend among workers. Carrying two devices is cumbersome, and many people simply preferred to use their new devices over corporate-issued phones or laptops. IT departments might have been able to brush this off, except that many of the early iPhone (and later, Android) adopters sat in executive offices. Now BYOD has spread around the world, creating a host of new challenges for IT departments concerning security, device management, and support costs.

More than half of information workers own the devices they use for work, according to Forrester Research, which surveyed almost 10 000 people in 17 countries, and that proportion is likely to increase, says David Johnson, a senior analyst at Forrester.

The groundswell caused many IT directors to simply throw up their hands. A study published last November by Kaspersky Lab, a digital-security firm, found that one in three organizations allowed personal cellphones unrestricted access to corporate resources—with troubling consequences. One in five companies in the same survey admitted losing business data after personal devices were lost or stolen.

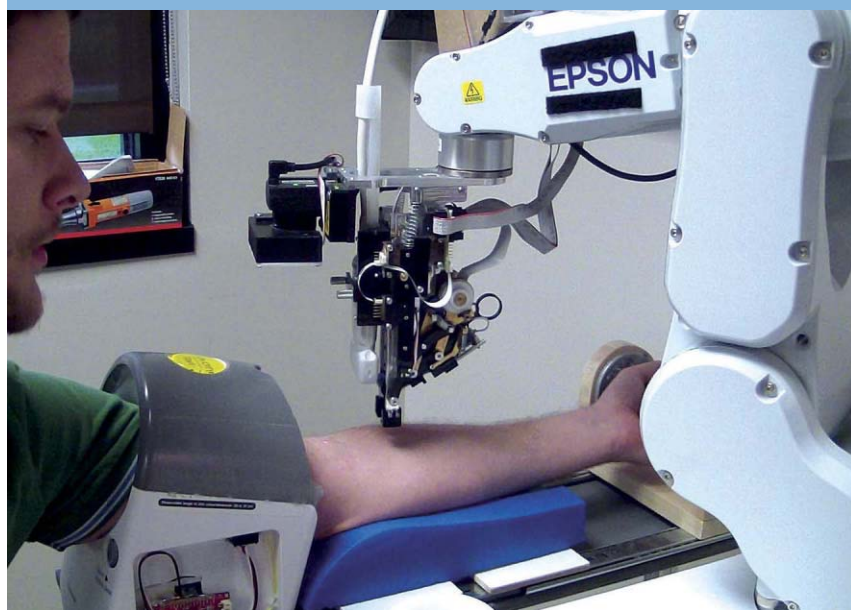
Among those companies that have tried to be more proactive, approaches vary widely. Some require the installation of software that allows the business to monitor all communications on a device and, if necessary, remotely wipe it. Others simply require acknowledgement of an acceptable use policy.

When Intel rolled out a BYOD program in 2010, the company's top concern was protecting its intellectual property, but it also

RESOURCES_START-UPS

PROFILE: VEEBOT

DRAWING BLOOD FASTER AND MORE SAFELY THAN A HUMAN CAN



ou probably know the routine for drawing blood. A medical technician briefly wraps your arm in a tourniquet and looks your veins over, some-

times tapping gently with a gloved finger on your inner elbow. Then the med tech selects a target. Usually, but not always, she gets a decent vein on the first try; sometimes it takes a second (or third) stick. This procedure is fine for the typical blood test at a doctor's office, but for contract researchers it represents a significant logistics problem. In drug trials it's not unusual to have to draw blood from dozens of people every hour or so throughout a day. These tests can add up to more than a hundred thousand blood draws a year for just one contract research company.

Veebot, a start-up in Mountain View, Calif., is hoping to automate drawing blood and inserting IVs by combining robotics with image-analysis software. To use the Veebot system, a patient puts his or her arm through an archway over a padded table. Inside the archway, an inflatable cuff

PLEASE HOLD STILL: Veebot's robot system can find a vein and place a needle at least as well as a human can. Clinical trials are expected to begin this year.

tightens around the arm, holding it in place and restricting blood flow to make the veins easier to see. An infrared light illuminates the inner elbow for a camera; software matches the camera's view against a model of vein anatomy and selects a likely vein. The vein is examined with ultrasound to confirm that it's large enough and has sufficient blood flowing through it. The robot then aligns the needle and sticks it in. The whole process takes about a minute, and the only thing the technician has to do is attach the appropriate test tube or IV bag.

Veebot began in 2009 when Richard Harris, a third-year undergraduate in Princeton's mechanical engineering department, was trying to come up with a topic for a project. At the same time, his father, Stuart Harris, founder of a company that does pharmaceutical contract research, mentioned that he'd love to see someone come up with a way to automate blood draws.

Harris says he was drawn to the idea because "it involved robotics and computer vision, both fields I was interested in, and it had demanding requirements because you'd be fully automating something that is different everytime and deals with humans."

He built a prototype that could find and puncture dots drawn on flexible plastic tubing, and with funding from his father, he cofounded Veebot in 2010.

Currently, Veebot's machine can correctly identify the best vein to target about 83 percent of the time, says Harris, which is about as good as a human. Harris wants to get that rate up to 90 percent before clinical trials. However, while he expects to achieve this in three to five months, he will then have to secure outside funding to cover the expense of those trials.

Harris estimates the market for his technology to be about US \$9 billion, noting that "blood is drawn a billion times a year in the U.S. alone; IVs are started 250 million times." Veebot will initially try to sell to large medical facilities.

Thomas Gunderson, managing director and a senior analyst at investment bank Piper Jaffray Companies, believes the time is right for this kind of medical device company. In a difficult case, "doctors today will search all over the hospital for the right person to do a blood draw, and they could still miss three or four times," he says. "Technology can help from a labor standpoint and make the procedure safer for the patient and for the person drawing the blood."

The biggest challenge, Harris says, is human psychology. "If people don't want a robot drawing their blood, then nobody is going to use it. We believe if this machine works better, faster, and cheaper than a person, people will want to use it."

Says Gunderson: "These days we have multimillion-dollar robots doing surgery. I think we passed 'creepy' several years ago and moved on." —TEKLA S. PERRY

Founded: 2010; **Headquarters:** Mountain View, Calif; **Founders:** Richard Harris, Stuart Harris, Joe Mygatt, James Wong; **Employees:** 4 **Funding:** undisclosed; **Website:** <http://www.veebot.com>

TECHNICALLY SPEAKING_BY PAUL MCFEDRIES

OPINION



TRACKING THE QUANTIFIED SELF

Self-tracking is not really a tool of optimization but of discovery, and if tracking regimes that we would once have thought bizarre are becoming normal, one of the most interesting effects may be to make us re-evaluate what “normal” means. —Gary Wolf, cofounder, The Quantified Self (The New York Times, 28 April 2010)

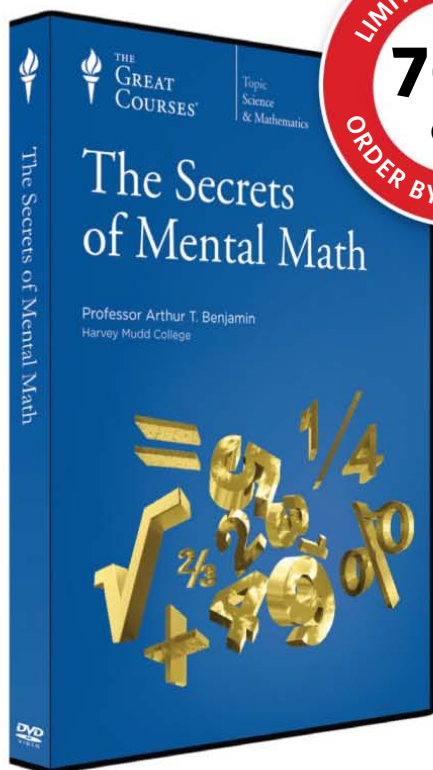


FOR BETTER OR FOR WORSE, we are data-generating machines. Whenever we pay with a credit card or drive through an automated toll system, or answer an e-mail or make a phone call, we can't help but leave a steady stream of ones and zeroes in our wake. This is our **digital exhaust**: the trackable or storable actions, choices, and preferences that we generate as we go about our daily lives. Even when just browsing the Web, we leave behind personal **clickprints** that uniquely identify our surfing behavior and lengthen the **paperless trails** that document our electronic selves. • From time to time, we harbor vague worries about oversharing on social networks or being tracked online by ad networks, but mostly we don't think about the **data shadows** that we cast wherever we go. But there is a growing segment of the population that spends a remarkable amount of time and effort trying to generate more personal data. While the rest of us are content to step on a bathroom scale once a week, these people weigh themselves several times a day. You and I might groggily estimate the number of hours we slept last night, but these people wear sensors that tell them how many hours and minutes of REM and non-REM sleep they achieved. • I speak, of course, of **self-trackers**, people who use technology to acquire, store, and analyze their own **life data**. Their **self-tracking** can also

create detailed records of food, exercise, location, and even mood, alertness, overall well-being, and other seemingly nonquantifiable psychological states. This process of **self-digitization** is often enhanced by smart clothing and other wearable computing technology that enables the **self-monitoring** of physiological states and **self-sensing** of such external data as location and time. These self-professed **data junkies** select from a variety of apps and websites that serve as tools for **self-quantifying**—and that prod them into doing even more of it. It is no wonder, then, that the movement as a whole is often called the quantified self (a term invented by *Wired* alums Gary Wolf and Kevin Kelly) and its practitioners are increasingly known as **quantified-selfers** or, simply, **QSers**.

You might think that the point of all this self-scrutiny is just to keep a record—that is, a **lifelog**—of vital stats, but self-trackers are not content with merely tracking a few numbers. Their interest lies in **quantitative assessment**: extracting *knowledge* from the raw data. They want to put their lives under the **macroscope**, which is the general term for any technology that enhances a person's ability to gather and analyze data. If that data tells you that you're just as bright-eyed and bushy-tailed on days when you managed only 5 or 6 hours of sleep, the lesson is clear: You're one of those lucky people who don't need 7 or 8 hours of sack time. If your heart rate and blood pressure spike when you sit down to dinner, maybe a little family counseling is in order. In short, by analyzing detailed data over a long period of time—either by generating charts in Excel or by using **auto-analytics** tools—self-trackers turn themselves into **self-experimenters**, perhaps even **body hackers**. The aim? Nothing more or less than the examined life, albeit one where “examined” means tracked, quantified, recorded, and analyzed.

But isn't all this *TMI* (too much information)—narcissism for gadget freaks and data geeks? Are these sorts of overexamined lives worth living? Sorry, I don't have enough data to answer those questions. ■



The Secrets of Mental Math

Taught by Professor Arthur T. Benjamin
HARVEY MUDD COLLEGE

LECTURE TITLES

1. Math in Your Head!
2. Mental Addition and Subtraction
3. Go Forth and Multiply
4. Divide and Conquer
5. The Art of Guesstimation
6. Mental Math and Paper
7. Intermediate Multiplication
8. The Speed of Vedic Division
9. Memorizing Numbers
10. Calendar Calculating
11. Advanced Multiplication
12. Masters of Mental Math

Discover the Secrets of Mental Math

One key to expanding your math potential—whether you're a corporate executive or a high-school student—lies in the power to perform mental math calculations. Solving basic math problems in your head offers lifelong benefits including a competitive edge at work, a more active and sharper mind, and improved performance on standardized tests.

In the 12 rewarding lectures of **The Secrets of Mental Math**, discover all the essential skills, tips, and tricks for improving and enhancing your ability to solve a range of basic math problems right in your head. Professor Arthur T. Benjamin, winner of numerous awards from the Mathematical Association of America, has designed this engaging course to be accessible to anyone looking to tap into his or her hidden mental calculating skills.

Offer expires 09/30/13

1-800-832-2412

WWW.THEGREATCOURSES.COM/3SPEC

The Secrets of Mental Math

Course no. 1406 | 12 lectures (30 minutes/lecture)

SAVE \$160

DVD ~~\$199.95~~ NOW \$39.95

+ \$5 Shipping, Processing, and Lifetime Satisfaction Guarantee
Priority Code: 78234

Designed to meet the demand for lifelong learning, The Great Courses is a highly popular series of audio and video lectures led by top professors and experts. Each of our more than 400 courses is an intellectually engaging experience that will change how you think about the world. Since 1990, over 10 million courses have been sold.

The COGNITIVE NET Is Coming

The Internet will break down without
new biologically inspired routing
By Antonio Liotta
Illustrations by L-Dopa

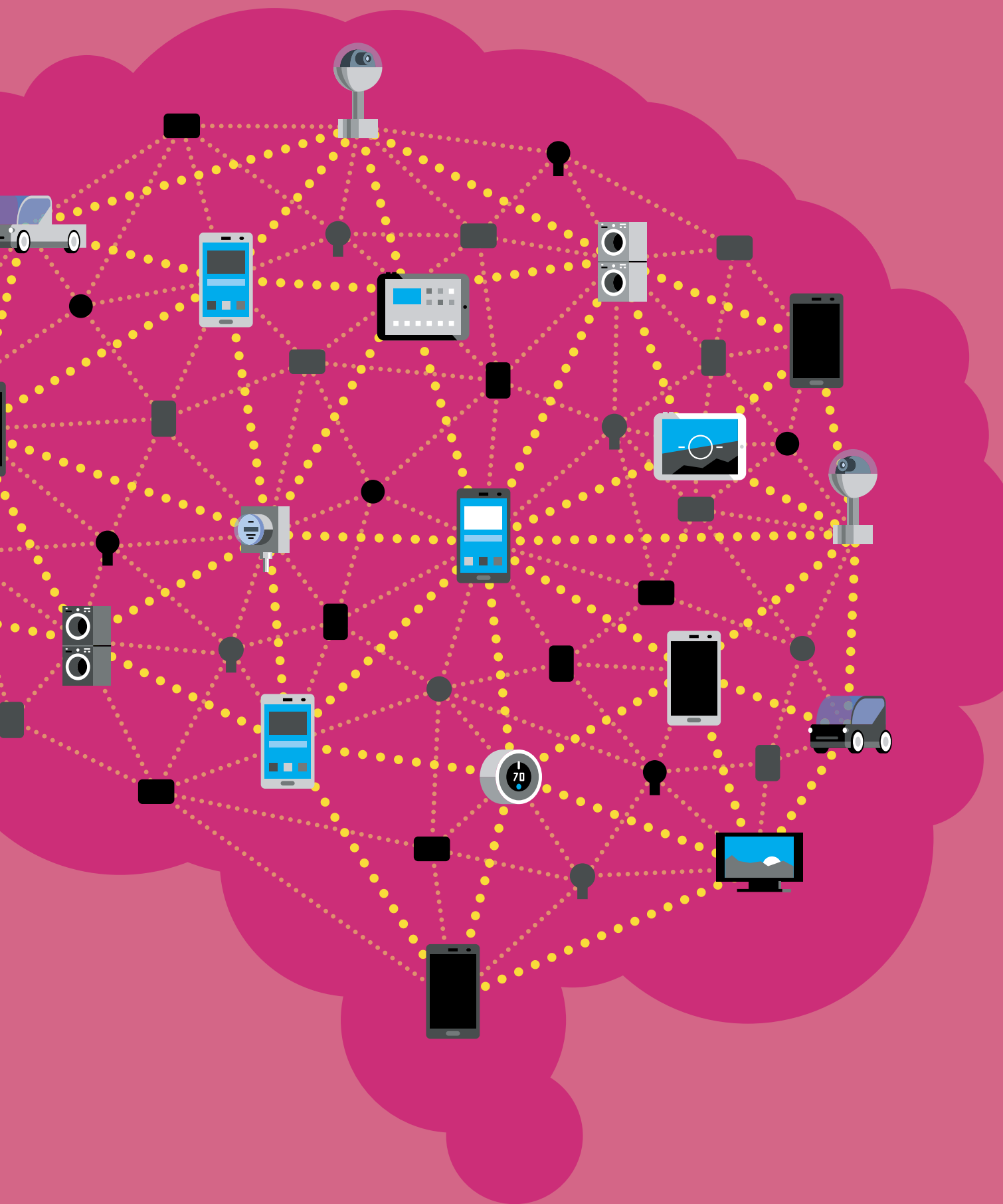
PERHAPS AS EARLY AS THE END OF THIS DECADE, our refrigerators will e-mail us grocery lists. Our doctors will update our prescriptions using data beamed from tiny monitors attached to our bodies. And our alarm clocks will tell our curtains when to open and our coffeemakers when to start the morning brew.

By 2020, according to forecasts from Cisco Systems, the global Internet will consist of 50 billion connected tags, televisions, cars, kitchen appliances, surveillance cameras, smartphones, utility meters, and whatnot. This is the Internet of Things, and what an idyllic concept it is.

But here's the harsh reality: Without a radical overhaul to its underpinnings, such a massive, variable network will likely create more problems than it proposes to solve. The reason? Today's Internet just isn't equipped to manage the kind of traffic that billions more nodes and diverse applications will surely bring.

In fact, it's already struggling to cope with the data being generated by ever-more-popular online activities, including video streaming, voice conferencing, and social gaming. Major Internet service providers around the world are now reporting global latencies greater than 120 milliseconds, which is about as much as a Voice over Internet Protocol connection can handle. Just imagine how slowly traffic would move if console gamers and cable television watchers, who now consume hundreds of exabytes of data off-line, suddenly migrated to cloud-based services.

The problem is not simply one of volume. Network operators will always be able to add capacity by transmitting data more efficiently and by rolling out more cables



and cellular base stations. But this approach is increasingly costly and ultimately unscalable, because the real trouble lies with the technology at the heart of the Internet: its routing architecture.

Information flows through the network using a four-decade-old scheme known as packet switching, in which data is sliced into small envelopes, or packets. Different packets may take different routes and arrive at different times, to be eventually reassembled at their destination. Routers, which decide the path each packet will take, are “dumb” by design. Ignorant of a packet’s origin and the bottlenecks it may encounter down the line, routers treat all packets the same way, regardless of whether they contain snippets of a video, a voice conversation, or an e-mail.

This arrangement worked superbly during the Internet’s early days. Back then most shared content, including e-mail and Web browsing, involved small sets of data transmitted with no particular urgency. It made sense for routers to process all packets equally because traffic patterns were mostly the same.

That picture has changed dramatically over the past decade. Network traffic today consists of bigger data sets, organized in more varied and complex ways. For instance, smart meters produce energy data in short, periodic bursts, while Internet Protocol television (IPTV) services generate large, steady streams. New traffic signatures will emerge as new applications come to market, including connected appliances and other products we haven’t yet imagined. Basic packet switching is just too rigid to manage such a dynamic load.

So it’s time we gave the Internet some smarts, not simply by making incremental improvements but by developing an entirely new way to transport data. And engineers are turning to nature for inspiration.

Millions of years of evolution have resulted in biological networks that have devised ingenious solutions to the hardest network problems, such as protecting against infectious agents and adapting to failures and changes. In particular, the human brain and body are excellent models for building better data networks. The challenge, of course, is in figuring out how to mimic them (see sidebar, “Networking Lessons From the Real World”).

To understand why the packet-switched Internet must be replaced with a more intelligent system, first consider how today’s network is structured. Say, for example, you want to watch a YouTube clip. For the video data to stream from Google’s server to your smartphone, the packets must pass through a hierarchy of subnetworks. They start at the outermost reaches of the Net: the access network, where terminals such as phones, sensors, servers, and PCs link up. Then the packets move through regional networks to the core network, or backbone. Here, dense fiber-optic cables ferry traffic at high speeds and across vast distances. Finally, the packets make their way back to the access network, where your smartphone resides.

Routers send each incoming packet along the best available route through this hierarchy. It works like this: Inside each router,

The Path to INTELLIGENT Routing

The future Internet will need smarter routing algorithms to handle diverse data flows and prevent failures. Although there are no tried-and-true solutions yet, early designs might follow an architecture like this one.

1 A ROUTING DEVICE can be any network node, such as a phone, a television, a car, a kitchen appliance, an environmental sensor, or some gadget yet to be invented. Proximal devices form “mesh networks” that off-load some traffic from the core network and bring Internet service to remote places.

2 THE ROUTING AND FORWARDING ENGINES determine the best pathways to get data packets to their destinations and queue them for transmission. (These engines are already built into today’s “dumb” routers, but in the future they could exist as software applications rather than separate pieces of hardware.)

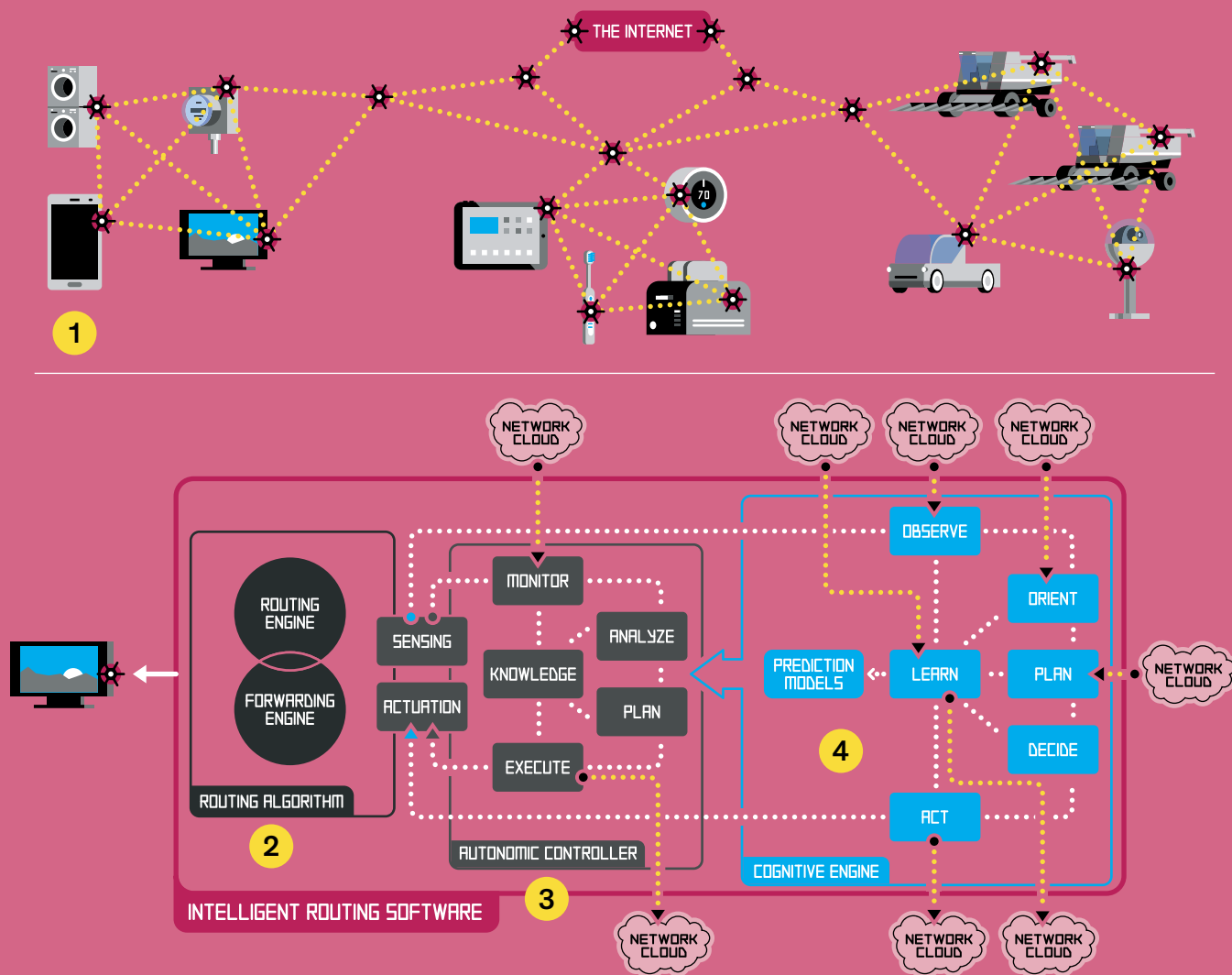
3 THE AUTONOMIC CONTROLLER directs the routing and forwarding engines by following the MAPE loop: It *monitors* internal sensor data and signals from other nodes, *analyzes* that information, *plans* a response, and *executes* it. Neighboring devices coordinate their actions in real time through control signals.

4 THE COGNITIVE ENGINE helps the router adapt to unforeseen changes by following the OOPDAL loop: It *observes* the environment, *orients* the system by prioritizing tasks, *plans* options, *decides* on a plan, *acts* on it, and *learns* from its actions. By sharing knowledge, devices spread intelligence across the Internet.

a collection of microchips called the routing engine maintains a table that lists the pathways to possible destinations. The routing engine continually updates this table using information from neighboring nodes, which monitor the network for signs of traffic jams. When a packet enters the router’s input port, another set of chips—the forwarding engine—reads the packet’s destination address and queries the routing table to determine the best node to send the packet to next. Then it switches the packet to a queue, or buffer, where it awaits transmission. The router repeats this process for each incoming packet.

There are several disadvantages to this design. First, it requires a lot of computational muscle. Table queries and packet buffering consume about 80 percent of a router’s CPU power and memory. And it’s slow. Imagine if a mail carrier had to recalculate the delivery route for each letter and package as it was collected. Routers likewise ignore the fact that many incoming packets may be headed for the same terminal.

Routers also overlook the type of data flow each packet belongs to. This is especially problematic during moments of peak traffic, when packets can quickly pile up in a router’s buffer. If more packets accumulate than the buffer can hold, the router discards



excess packets somewhat randomly. In this scenario, a video stream—despite having strict delivery deadlines—would experience the same packet delays and losses as an e-mail. Similarly, a large file transfer could clog up voice and browsing traffic so that no single flow reaches its destination in a timely manner.

And what happens when a crucial routing node fails, such as when a Vodafone network center in Rotterdam, Netherlands, caught fire in 2012? Ideally, other routers will figure out how to divert traffic around the outage. But often, local detours just move the congestion elsewhere. Some routers become overloaded with packets, causing more rerouting and triggering a cascade of failures that can take down large chunks of the network. After the Vodafone fire, 700 mobile base stations were out of commission for more than a week.

Routers could manage data flows more effectively if they made smarter choices about which packets to discard and which ones to expedite. To do this, they would need to gather much more information about the network than simply the availability of routing links. For instance, if a router knew it was receiving high-quality IPTV packets destined for a satellite phone, it might choose to drop those packets in order to prioritize others that are more likely to reach their destinations.

Ultimately, routers will have to coordinate their decisions and actions across all levels of the Internet, from the backbone to the end terminals, and the applications running on them. And as new user devices, services, and threats come on line in the future, the system will need to be smart enough to adapt.

The first step in designing a more intelligent Internet is to endow every connected computer with the ability to route data. Given the computational capabilities of today's consumer devices, there's no reason for neighboring smart gadgets to communicate over the core network. They could instead use any available wireless technology, such as Wi-Fi or Bluetooth, to spontaneously form "mesh networks." This would make it possible for any terminal that taps into the access network—tablet, television, thermostat, tractor, toaster, toothbrush, you name it—to relay data packets on behalf of any other terminal.

By off-loading local traffic from the Internet, mesh networks would free up bandwidth for long-distance services, such as IPTV, that would otherwise require costly infrastructure upgrades. These networks would also add routing pathways that bypass bottlenecks, so traffic could flow to areas where Internet access is now

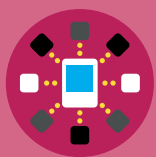
Networking LESSONS From the Real World

Internet engineers can learn a lot from biological and social networks



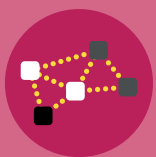
Keep pathways short, even in large networks.

EXAMPLES: Social relationships, gene regulation, neural networks in the brain
ADVANTAGES: When data can reach any destination in a small number of steps, latency stays low.



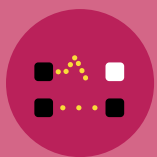
Only a small percentage of nodes should have many links.

EXAMPLES: Human sexual partners, scientific-paper citations
ADVANTAGES: Minimizing the number of hubs helps stop the spread of viruses and protects against attacks.



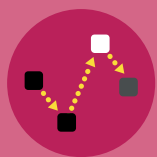
Weak links can be a good thing.

EXAMPLES: Some molecular structures
ADVANTAGES: Poor or transient links can help improve bad connections, dissipate disruptions, and bring network access to places where strong links can't be built.



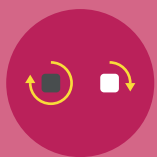
Trade some speed for stability.

EXAMPLES: Traffic-control systems (including stoplights, yield signs, and speed limits)
ADVANTAGES: Controlling data flows helps prevent traffic spikes from causing network congestion or collapse.



Spread information through "gossip" rather than broadcast.

EXAMPLES: Rumors, viral videos
ADVANTAGES: Disseminating data as if it were gossip can be more efficient and less disruptive than broadcasting it.



Control and learn at different timescales.

EXAMPLES: Autonomic functions (such as breathing and digesting) versus cognition
ADVANTAGES: Real-time control lets nodes coordinate actions, while gradual learning helps the network evolve.

poor, extending cellular service underground, for example, and providing extra coverage during natural disasters.

But to handle data and terminals of many different kinds, routers (including the terminals themselves) need better methods for building and selecting data pathways. One way to engineer these protocols is to borrow tricks from a complex network that already exists in nature: the human autonomic nervous system.

This system controls breathing, digestion, blood circulation, body heat, the killing of pathogens, and many other bodily functions. It does all of this, as the name suggests, autonomously—without our direction or even our awareness. Most crucially, the autonomic nervous system can detect disturbances and make adjustments before these disruptions turn into life-threatening problems.

If all this sounds a little vague, consider the example of digestion. Say you've just eaten a big, juicy hamburger. To begin breaking it down, the stomach must secrete the proper amount of gastric juices. This might seem like a simple calculation: more meat, more juices. In fact, the parts of the brain that control this process rely on a smorgasbord of inputs from many other systems, including taste, smell, memory, blood flow, hormone levels, muscle activity, and immune responses. Does that burger contain harmful bacteria that must be killed or purged? Does the body need to conserve blood and fuel for more important tasks, such as running from an enemy? By coordinating many different organs and functions at once, the autonomic system keeps the body running smoothly.

By contrast, the Internet addresses a disturbance, such as a spike in traffic or a failed node, only after it starts causing trouble. Routers, servers, and computer terminals all try to fix the problem separately, rather than work together. This often just makes the problem worse—as was the case during the Vodafone fire.

A more cooperative Internet requires routing and forwarding protocols that behave more like the autonomic nervous system. Network engineers are still figuring out how best to design such a system, and their solutions will no doubt become more sophisticated as they work more closely with biologists and neuroscientists.

One idea, proposed by IBM, is the Monitor-Analyze-Plan-Execute (MAPE) loop, or more simply, the knowledge cycle. Algorithms that follow this architecture must perform four main tasks:

First, they *monitor* a router's environment, such as its battery level, its memory capacity, the type of traffic it's seeing, the number of nodes it's connected to, and the bandwidth of those connections.

Then the knowledge algorithms *analyze* all that data. They use statistical techniques to determine whether the inputs are typical and, if they aren't, whether the router can handle them. For example, if a router that typically receives low-quality video streams suddenly receives a high-quality one, the algorithms calculate whether the router can process the stream before the video packets fill its buffer.

Next, they *plan* a response to any potential problem, such as an incoming video stream that's too large. For instance, they may figure the best plan is to ask the video server to lower the stream's bit rate. Or they may find it's better to break up the stream and work with other nodes to spread the data over many different pathways.

Lastly, they *execute* the plan. The execution commands may modify the routing tables, tweak the queuing methods, reduce transmission power, or select a different transmission channel, among many possible actions.

A routing architecture like the MAPE loop will be key to keeping the Internet in check. Not only will it help prevent individual routers from failing, but by monitoring data from neighboring

nodes and relaying commands, it will also create feedback loops within the local network. In turn, these local loops swap information with other local networks, thereby propagating useful intelligence across the Net.

It's important to note that there's no magic set of algorithms that will work for every node and every local network. Mesh networks of smartphones, for example, may operate best using protocols based on swarm intelligence, such as the system ants use to point fellow ants to a food source. Meanwhile, massive monitoring networks, such as "smart dust" systems made of billions of grain-size sensors, may share data much as people share gossip—a method that would minimize transmission power.

Autonomic protocols would help the Internet better manage today's traffic flows. But because new online services and applications emerge over the lifetime of any router, routers will have to be able to learn and evolve on their own.

To make this happen, engineers must turn to the most evolutionarily advanced system we know: human cognition. Unlike autonomic systems, which rely on predetermined rules, cognitive systems make decisions based on experience. When you reach for a ball flying toward you, for example, you decide where to position your hand by recalling previous successes. If you catch the ball, the experience reinforces your reasoning. If you drop the ball, you'll revise your strategy.

Of course, scientists don't know nearly enough about natural cognition to mimic it exactly. But advances in the field of machine learning—including pattern-recognition algorithms, statistical inference, and trial-and-error learning techniques—are proving to be useful tools for network engineers. With these tools, it's possible to create an Internet that can learn to juggle unfamiliar data flows or fight new malware attacks in a manner similar to the way a single computer might learn to recognize junk mail or play "Jeopardy!"

Engineers have yet to find the best framework for designing cognitive networks. A good place to start, however, is with a model first proposed in the late 1990s for building smart radios. This architecture is known as the cognition cycle, or the Observe-Orient-Plan-Decide-Act-Learn (OOPDAL) loop. Like the MAPE loop in an autonomic system, it begins with the *observation* of environmental conditions, including internal sensor data and signals from nearby nodes. Cognition algorithms then *orient* the system by evaluating and prioritizing the gathered information. Here things get more complex. For low-priority actions, the algorithms consider alternative *plans*. Then they *decide* on a plan and *act* on it, either by triggering new internal behavior or by signaling nearby nodes. When more-urgent action is needed, the algorithms can bypass one or both of the planning and decision-making steps. Finally, by observing the results of these actions, the algorithms *learn*.

In an Internet router, OOPDAL loops would run parallel to the autonomic MAPE loop (see illustration, "The Path to Intelligent Routing"). As the cognition algorithms learned, they would generate prediction models that would continually modify the knowledge algorithms, thereby improving the router's ability to manage diverse data flows. This interaction is akin to the way

your conscious brain might retrain your arm muscles to catch a hardball after years of playing with a softball.

Network engineers are still far from creating completely cognitive networks, even in the laboratory. One of the biggest challenges is designing algorithms that can learn not only how to minimize the use of resources—such as processing power, memory, and radio spectrum—but also how to maximize the quality of a user's experience. This is no trivial task. After all, experience can be highly subjective. A grainy videoconference might be a satisfactory experience for a teenager on a smartphone, but it would be unacceptable to a business executive chatting up potential clients. Likewise, you might be more tolerant of temporary video freezes if you were watching a free television service than if you were paying for a premium plan.

Nevertheless, my colleagues and I at the Eindhoven University of Technology, in the Netherlands, have made some progress. Using a network emulator, or "Internet in a box," we can simulate various network conditions and test how they affect the perceived quality of different types of video streams. In our experiments, we have identified hundreds of measurable parameters to predict the quality of experience, including latency, jitter, video content, image resolution, and frame rate. Using new sensing protocols, terminals could also measure things like the type of screen someone's using, the distance between the screen and the user, and the lighting conditions in the room.

In collaboration with Telefónica, in Spain, we have created machine-learning algorithms that use many of these parameters to predict the quality of a user's experience when IPTV programs are streamed to different types of smartphones. These prediction models turned out to be remarkably accurate (having around a 90 percent agreement with user surveys), showing that it's possible to train networks to adapt to variable conditions on their own. In another study, we demonstrated that a network can quickly learn, through trial and error, the best bit rate for delivering a specific video stream with the highest possible quality of experience. One big advantage of this method is that it can be applied to any type of network and any type of video, whether the network has seen it before or not.

Engineers still have plenty of work to do before they can build complex intelligence into the Internet itself. Although the change won't happen overnight, it's already beginning. At the edges of the network, services such as Google and Facebook are now using sophisticated learning algorithms to infer our preferences, make recommendations, and customize advertisements. Wireless equipment manufacturers are building radios that can select frequencies and adjust their transmission power by "listening" to the airwaves. Still other engineers are finalizing protocols for creating mobile ad hoc networks so that police and rescue vehicles, for example, can communicate directly with one another.

Gradually, similar innovations will spread to other parts of the network. Perhaps as early as 2030, large portions of the Internet could be autonomic, while others will show the odd flash of actual insight. The future Net will exhibit a great diversity of intelligence, much like our planet's own biological ecosystems. ■

POST YOUR COMMENTS online at <http://spectrum.ieee.org/cognitive0813>

CLAD IN CONTROVERSY

AFTER YEARS OF RESEARCH,
DEVELOPMENT, AND DEBATE,
THE USS *ZUMWALT*, THE FIRST OF
A NEW CLASS OF HIGH-TECH
DESTROYERS, NEARS COMPLETION

BY DAVID SCHNEIDER

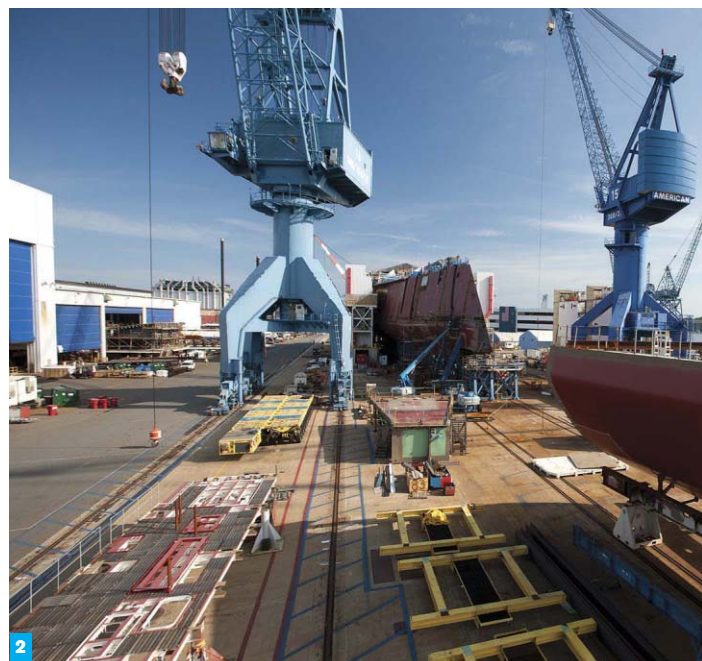
Two decades ago, the U.S. Navy began designing what it then called its “21st-century destroyers.” These were to be a fleet of 32 guided-missile destroyers that would be able to cruise near coastlines and attack forces on land with mind-boggling might. In 2001, though, the Navy canceled that program and replaced it with a less costly alternative.

It took another dozen years, but the first destroyer of that new generation is now nearing completion. Although less ambitious than the original concept, the first ship of this new class, the USS *Zumwalt*, is pioneering so many advanced technologies that some decision makers have criticized the program for trying to do too much, too soon.

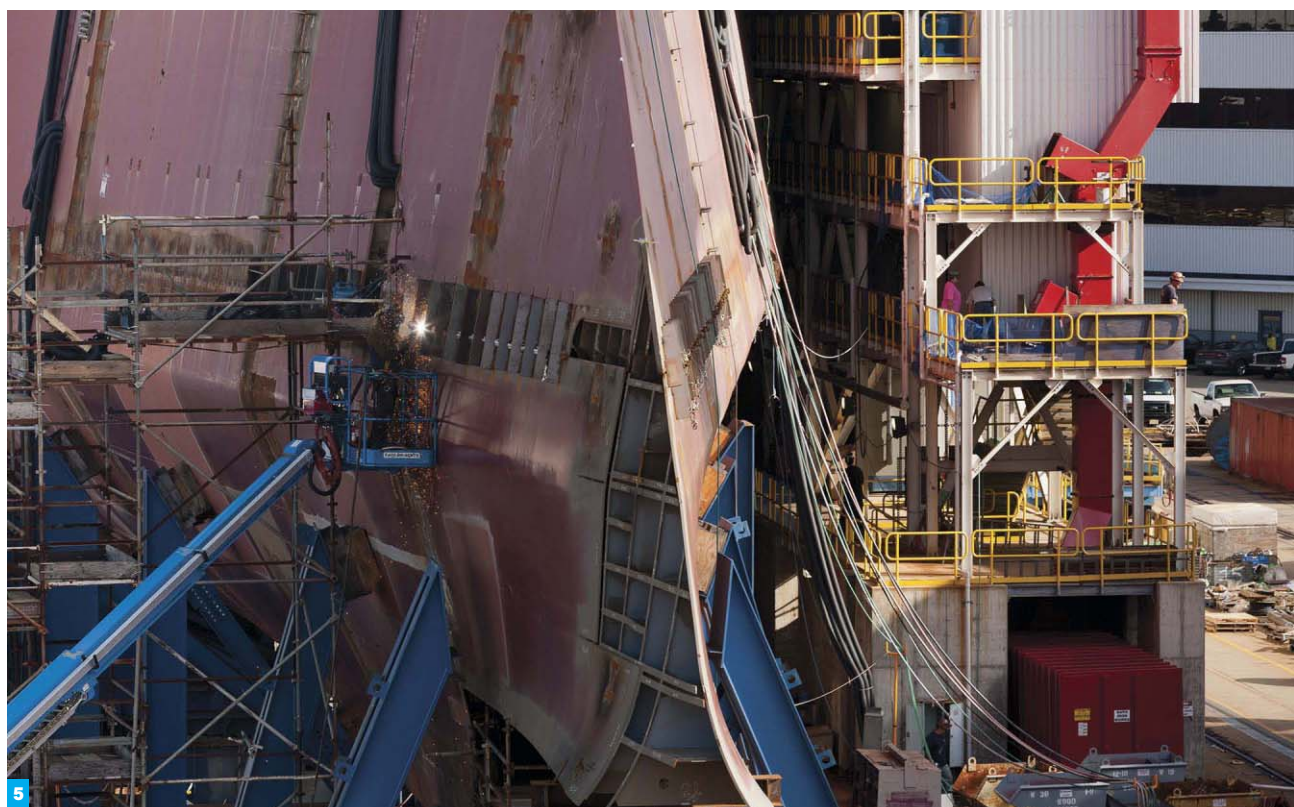
Some of the pushback came simply from the enormous costs involved in developing so many cutting-edge technologies. Indeed, faced with mushrooming costs, the U.S. government reacted by repeatedly reducing the number of these destroyers to be built—eventually settling on just three ships. The total cost of the program, including R&D, that will result in those three ships is estimated to be US \$22 billion, according to the Congressional Research Service. Another point of intense debate was whether the main task

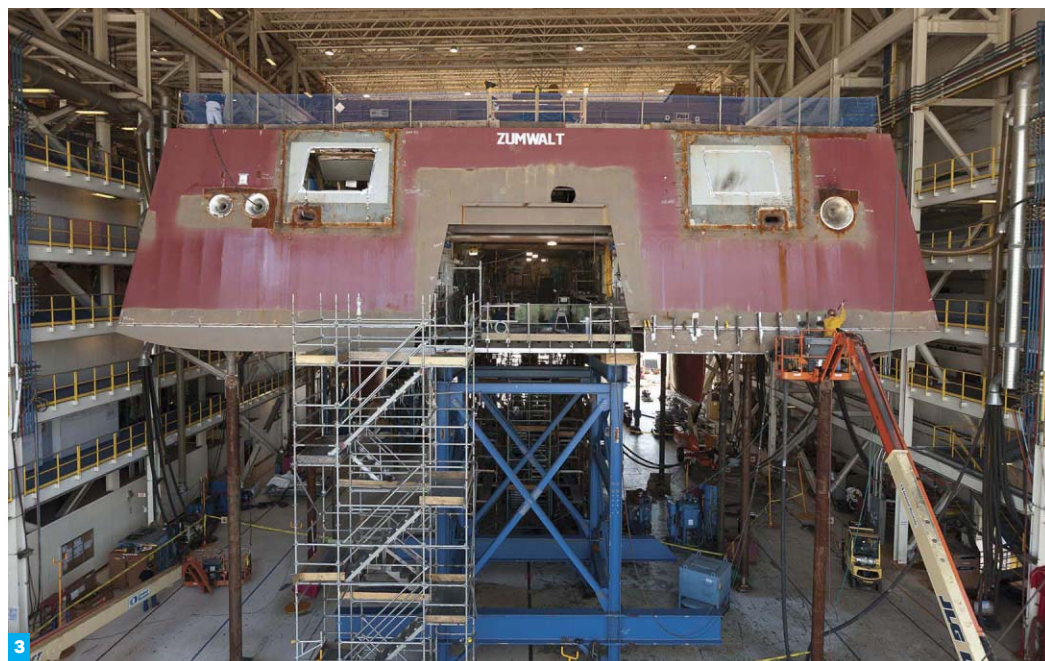






CLOCKWISE FROM TOP LEFT: 1. The 24-meter beam of this ship is evident in this view, showing a cross section from just behind the aircraft hangar. **2.** The 186-meter-long ship was assembled in huge sections at Bath Iron Works, in Maine. **3.** Some of those sections were fabricated within this shipyard's indoor facilities before being transferred outside. **4.** The upper part of the ship's deckhouse, made of balsawood-cored carbon-composite panels, is comparatively light. **5.** Welders work on the oddly shaped bow section.





envisioned for this ship—cruising in coastal waters while supporting military operations on nearby lands—was really so important to U.S. geopolitical interests.

And there has been no shortage of purely technical questions. Chief among them: Are the many advanced technologies slated for the *Zumwalt* really battle ready? It will

probably be years before we'll know for sure. But it's not too soon to consider how these technologies will affect future naval warfare.

The U.S. Navy has not released details about the ship's interior. But after gathering what information we could, including construction photos, we assembled the accompanying illustration. Together these visual

elements offer what may be a preview of how warships will look for decades to come.

ONE OF THE MOST OBVIOUS differences between the *Zumwalt* and almost all other ships is its basic shape. The *Zumwalt* has what's known as a tumblehome hull, which narrows rather than widens with height above

COMPOSITE DECKHOUSE

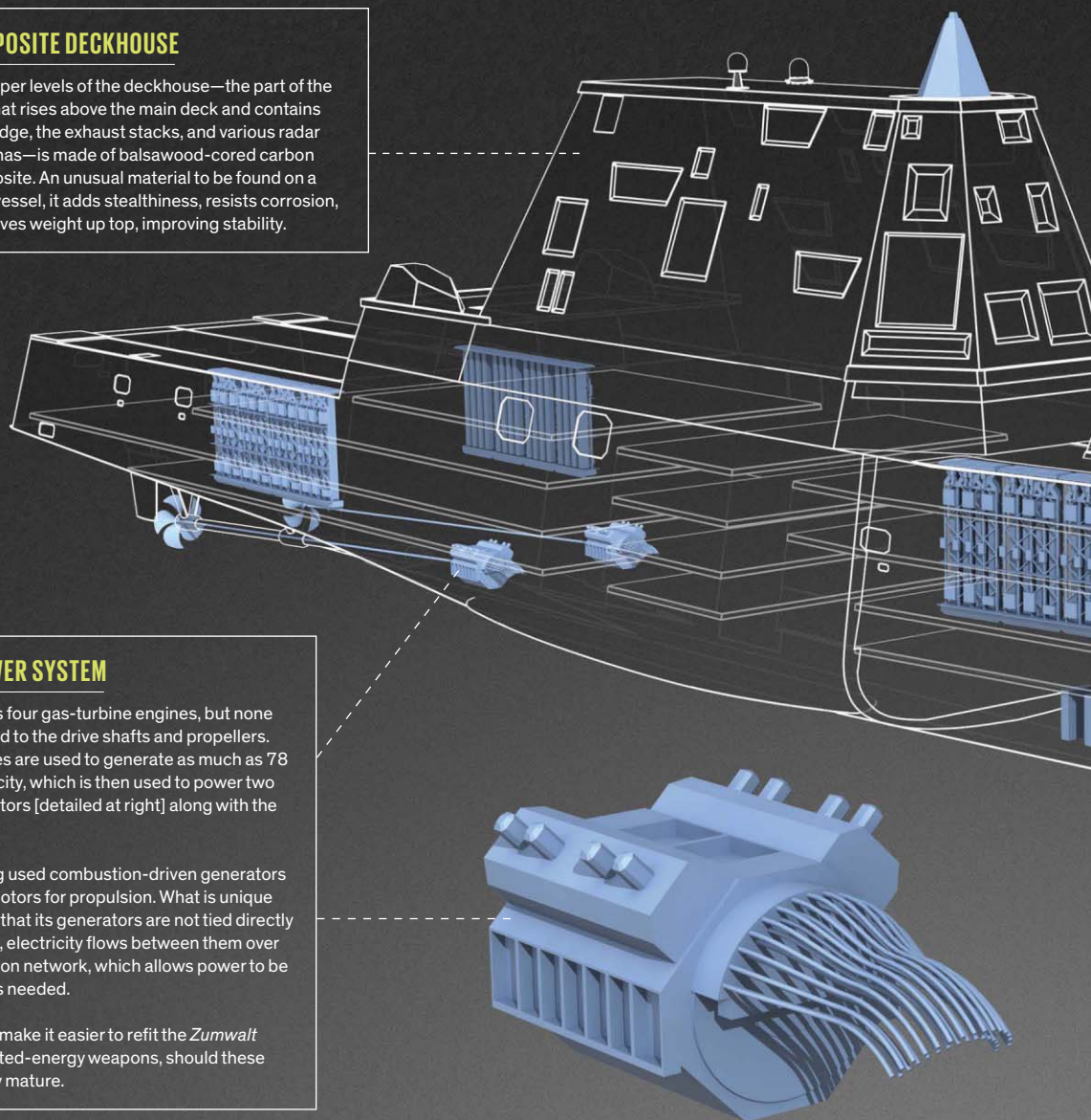
The upper levels of the deckhouse—the part of the ship that rises above the main deck and contains the bridge, the exhaust stacks, and various radar antennas—is made of balsawood-cored carbon composite. An unusual material to be found on a naval vessel, it adds stealthiness, resists corrosion, and saves weight up top, improving stability.

INTEGRATED POWER SYSTEM

The USS *Zumwalt* has four gas-turbine engines, but none are directly connected to the drive shafts and propellers. Instead, these engines are used to generate as much as 78 megawatts of electricity, which is then used to power two electric induction motors [detailed at right] along with the ship's other systems.

Large ships have long used combustion-driven generators linked with electric motors for propulsion. What is unique about the *Zumwalt* is that its generators are not tied directly to its motors. Instead, electricity flows between them over a ship-wide distribution network, which allows power to be directed wherever it's needed.

This flexibility should make it easier to refit the *Zumwalt* with railguns or directed-energy weapons, should these technologies one day mature.



the waterline. The rake of the bow is also inverted, making the ship look like an oddly angular submarine.

Tumblehome hulls haven't been seen on naval ships in over a century. The U.S. Navy used it here because the inward-angled hull won't reflect radar energy straight back to an adversary's antennas. Its main disadvantage is instability: A tumblehome hull provides no additional righting force when the ship heels over, causing some naval architects to speculate that it could make the *Zumwalt* prone to capsizing in rough seas.

Another distinguishing feature of the *Zumwalt* is its deckhouse, which rises above the main deck and houses the bridge, the

exhaust stacks, and various radar antennas. Like the hull, it was designed to reduce the ship's radar profile and has sides that cant inward. Unlike the steel hull, the upper part of the deckhouse is made of balsawood-cored carbon-composite panels.

This material, highly unusual for a warship, was used to reduce weight up top (which aids stability), enhance corrosion resistance, and make the ship more stealthy. But it's very expensive, and in January of this year the Navy began investigating using only steel for the deckhouse of the third and final ship of the *Zumwalt* class, the USS *Lyndon B. Johnson*.

Yet another departure from tradition is how the *Zumwalt* arranges its many missiles.

Guided-missile destroyers of earlier design position their vertical missile-launcher tubes amidships, where they are best protected from enemy fire. The *Zumwalt*'s designers arrayed its missiles along the ship's flanks, positioning them between inner and outer hulls. Putting them on the periphery does make the missiles more vulnerable to enemy fire, but it lessens the consequences should they be struck. Were that to happen, the resulting blast would explode outward, leaving intact the inner, watertight hull.

In another break from the U.S. Navy's usual designs, the *Zumwalt*'s propellers and drive shafts are turned by electric motors, rather than being directly attached to combustion

PERIPHERAL VERTICAL LAUNCH SYSTEM

On other destroyers, missiles are stored amidships, so as to be best protected from enemy fire. On the *Zumwalt*, missiles are arrayed around the ship's exterior. The vertical missile-launcher cells, sandwiched between inner and outer hulls, are designed to explode outward if they are hit, leaving the watertight inner hull intact.

ADVANCED GUN SYSTEM

The ship's two 155-millimeter guns fire self-propelled projectiles that can be guided in flight. They are capable of reaching targets more than 100 kilometers away. Each gun can hold in excess of 300 of these high-tech rounds, which are handled by an automated loading mechanism.

TUMBLEHOME HULL

As the hull rises from the waterline, its side angles inward, a feature not seen in naval warships in more than a century. It is used here to reduce the ship's radar profile.

engines. Such electric-drive systems, while a rarity for the U.S. Navy, have long been standard on big ships. What's new and different about the one on the *Zumwalt* is that it's flexible enough to propel the ship, fire railguns or directed-energy weapons (should these eventually be deployed), or both at the same time. That's because the 78 megawatts from its four gas-turbine generators can be directed through the ship's power-distribution network wherever it's needed. The presence of such a tightly integrated power-generation and distribution system has led some to call the *Zumwalt* the U.S. Navy's first "all-electric ship."

While the general idea of using electric motors to propel the ship wasn't particularly

controversial, the choice of what kind of motors to use did not come easily. The leading idea at first was to use permanent-magnet motors, but these proved challenging to develop, and the Navy ultimately opted for two 34-MW induction motors instead.

IT'S PERHAPS A BIT IRONIC that, despite the many cutting-edge technologies it contains, the *Zumwalt* class was passed over for one of the Navy's most technologically challenging missions of all: sea-based ballistic-missile defense, which has grown more important to the United States and its allies lately as more nations of concern attain nuclear and ballistic-missile capa-

bilities. Instead, the Navy will build more destroyers of a more conventional type and outfit them with the radars and anti-ballistic missiles needed.

In a 2009 speech, Adm. Gary Roughhead, then Chief of Naval Operations, made his reasoning for this change clear. While he applauded the *Zumwalt*'s advanced technology and how the program was being run, he also repeated a truism that only the most naive engineers in attendance didn't already know: "Technology does not always equate to relevant capability." ■

POST YOUR COMMENTS online at <http://spectrum.ieee.org/electricship0813>



ONLY CONNECT: Researcher Hubert Zimmermann [left] explains computer networking to French officials at a meeting in 1974. Zimmermann would later play a key role in the development of the Open Systems Interconnection standards.

THE INTERNET THAT WASN'T

The making—and forgetting—of the
Open Systems Interconnection standards

BY ANDREW L. RUSSELL

If everything had gone according to plan, the Internet as we know it would never have sprung up. That plan, devised 35 years ago, instead would have created a comprehensive set of standards for computer networks called Open Systems Interconnection, or OSI. Its architects were a dedicated group of computer industry representatives in the United Kingdom, France, and the United States who envisioned a complete, open, and multilayered system that would allow users all over the world to exchange data easily and thereby unleash new possibilities for collaboration and commerce.

For a time, their vision seemed like the right one. Thousands of engineers and policy-makers around the world became involved in the effort to establish OSI standards. They soon had the support of everyone who mattered: computer companies, telephone companies, regulators, national governments, international standards setting agencies, academic researchers, even the U.S. Department of Defense. By the mid-1980s the worldwide adoption of OSI appeared inevitable.

And yet, by the early 1990s, the project had all but stalled in the face of a cheap and agile, if less comprehensive, alternative: the Internet's Transmission Control Protocol and Internet Protocol. As OSI faltered, one of the Internet's chief advocates, Einar Stefferud, gleefully pronounced: "OSI is a beautiful dream, and TCP/IP is living it!"

What happened to the "beautiful dream"? While the Internet's triumphant story has been well documented by its designers and the historians they have worked with, OSI has been forgotten by all but a handful of veterans of the Internet-OSI standards wars. To understand why, we need to dive into the early history of computer networking, a time

INFRA

when the vexing problems of digital convergence and global interconnection were very much on the minds of computer scientists, telecom engineers, policymakers, and industry executives. And to appreciate that history, you'll have to set aside for a few minutes what you already know about the Internet. Try to imagine, if you can, that the Internet never existed.

The story starts in the 1960s. The Berlin Wall was going up. The Free Speech movement was blossoming in Berkeley. U.S. troops were fighting in Vietnam. And digital computer-communication systems were in their infancy and the subject of intense, wide-ranging investigations, with dozens (and soon hundreds) of people in academia, industry, and government pursuing major research programs.

The most promising of these involved a new approach to data communication called packet switching. Invented independently by Paul Baran at the Rand Corp. in the United States and Donald Davies at the National Physical Laboratory in England, packet switching broke messages into discrete blocks, or packets, that could be routed separately across a network's various channels. A computer at the receiving end would reassemble the packets into their original form. Baran and Davies both believed that packet switching could be more robust and efficient than circuit switching, the old technology used in telephone systems that required a dedicated channel for each conversation.

Researchers sponsored by the U.S. Department of Defense's Advanced Research Projects Agency created the first packet-switched network, called the ARPANET, in 1969. Soon other institutions, most notably the computer giant IBM and several of the telephone monopolies in Europe, hatched their own ambitious plans for packet-switched networks. Even as these institutions contemplated the digital convergence of computing and communications, however, they were anxious to protect the revenues generated by their existing businesses. As a result, IBM and the telephone monopolies favored packet switching that relied on "virtual circuits"—a design that mimicked circuit switching's technical and organizational routines.

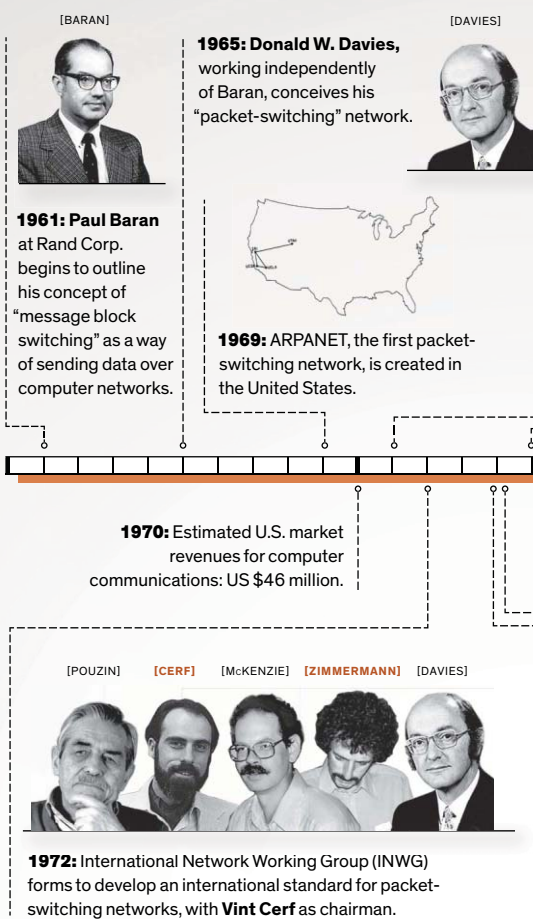
With so many interested parties putting forth ideas, there was widespread agree-

ment that some form of international standardization would be necessary for packet switching to be viable. An early attempt began in 1972, with the formation of the International Network Working Group (INWG). Vint Cerf was its first chairman; other active members included Alex McKenzie in the United States, Donald Davies and Roger Scantlebury in England, and Louis Pouzin and Hubert Zimmermann in France.

The purpose of INWG was to promote the "datagram" style of packet switching that Pouzin had designed. As he explained to me when we met in Paris in 2012, "The essence of datagram is connectionless. That means you have no relationship established between sender and receiver. Things just go separately, one by one, like photons." It was a radical proposal, especially when compared to the connection-oriented virtual circuits favored by IBM and the telecom engineers.

INWG met regularly and exchanged technical papers in an effort to reconcile its designs for datagram networks, in particular for a transport protocol—the key mechanism for exchanging packets across different types of networks. After several years of debate and discussion, the group finally reached an agreement in 1975, and Cerf and Pouzin submitted their protocol to the international body responsible for overseeing telecommunication standards, the International Telegraph and Telephone Consultative Committee (known by its French acronym, CCITT).

The committee, dominated by telecom engineers, rejected the INWG's proposal as too risky and untested. Cerf and his colleagues were bitterly disappointed. Pouzin, the combative leader of Cyclades, France's own packet-switching research project, sarcastically noted that members of the CCITT "do not object to packet switching, as long as it looks just like circuit switching." And when Pouzin complained at major conferences about the "arm-twisting" tactics of "national monopolies," everyone knew he was referring to the French telecom authority. French bureaucrats did not appreciate their countryman's candor, and government funding was



drained from Cyclades between 1975 and 1978, when Pouzin's involvement also ended.

For his part, Cerf was so discouraged by his international adventures in standards making that he resigned his position as INWG chair in late 1975. He also quit the faculty at Stanford and accepted an offer to work with Bob Kahn at ARPA. Cerf and Kahn had already drawn on Pouzin's datagram design and published the details of their "transmission control program" the previous year in the *IEEE Transactions on Communications*. That provided the technical foundation of the "Internet," a term adopted later to refer to a network of networks that utilized ARPA's TCP/IP. In subsequent years the two men directed the development of Internet protocols in an environment they could control: the small community of ARPA contractors.

Cerf's departure marked a rift within the INWG. While Cerf and other ARPA con-

A BRIEF HISTORY OF THE OSI STANDARDS

1971: Cyclades packet-switching project launches in France.

1975: INWG submits a proposal to the International Telegraph and Telephone Consultative Committee (CCITT), which rejects it. Cerf resigns from INWG.

[BACHMAN] [ZIMMERMANN] [DAY]



1977: International Organization for Standardization (ISO) committee on Open Systems Interconnection is formed, with **Charles Bachman** as chairman.

January **1983:** U.S. Department of Defense's mandated use of TCP/IP on the ARPANET signals the "birth of the Internet."

May **1983:** ISO publishes "ISO 7498: The Basic Reference Model for Open Systems Interconnection" as an international standard.

1985: U.S. National Research Council recommends that the Department of Defense migrate gradually from TCP/IP to OSI.

1989: As OSI begins to founder, computer scientist Brian Carpenter gives a talk entitled "Is OSI Too Late?" He receives a standing ovation.

WWW

1991: Tim Berners-Lee announces public release of the WorldWideWeb application.

1974: IBM launches a packet-switching network called the Systems Network Architecture.

[CERF]



[KAHN]

1974: Cerf and **Robert Kahn** publish "A Protocol for Packet Network Intercommunication," in *IEEE Transactions on Communications*.

1980: U.S. Department of Defense publishes "Standards for the Internet Protocol and Transmission Control Protocol."

1976: CCITT publishes Recommendation X.25, a standard for packet switching that uses "virtual circuits."

1988: U.S. market revenues for computer communications: \$4.9 billion.



1988: U.S. Department of Commerce mandates that government agencies buy OSI-compliant products.



1992: U.S. National Science Foundation revises policies to allow commercial traffic over the Internet.

1992: In a "palace revolt," Internet engineers reject the ISO ConnectionLess Network Protocol as a replacement for IP version 4.

IPv6

1996: Internet community defines IP version 6.

tractors eventually formed the core of the Internet community in the 1980s, many of the remaining veterans of INWG regrouped and joined the international alliance taking shape under the banner of OSI. The two camps became bitter rivals.

OSI was devised by committee, but that fact alone wasn't enough to doom the project—after all, plenty of successful standards start out that way. Still, it is worth noting for what came later.

In 1977, representatives from the British computer industry proposed the creation of a new standards committee devoted to packet-switching networks within the International Organization for Standardization (ISO), an independent nongovernmental association created after World War II. Unlike the CCITT, ISO wasn't specifically concerned with telecommunications—the wide-ranging topics of its technical committees included

TC1 for standards on screw threads and TC17 for steel. Also unlike the CCITT, ISO already had committees for computer standards and seemed far more likely to be receptive to connectionless datagrams.

The British proposal, which had the support of U.S. and French representatives, called for "network standards needed for open working." These standards would, the British argued, provide an alternative to traditional computing's "self-contained, 'closed' systems," which were designed with "little regard for the possibility of their interworking with each other." The concept of open working was as much strategic as it was technical, signaling their desire to enable competition with the big incumbents—namely, IBM and the telecom monopolies.

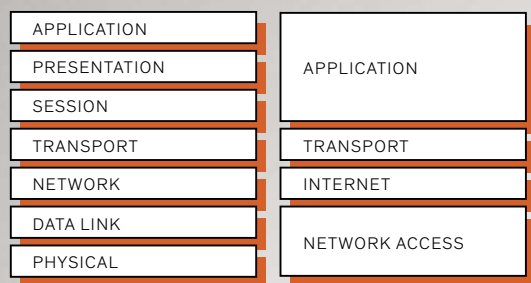
As expected, ISO approved the British request and named the U.S. database expert Charles Bachman as committee chairman. Widely respected in computer circles,

BARAN: COMPUTER HISTORY MUSEUM; DAVIES: NPL; MAP AND BACHMAN: COMPUTER HISTORY MUSEUM; ZIMMERMANN AND DAY: JOHN DAY; POUZIN: MARC WEBER/COMPUTER HISTORY MUSEUM; CERF: JOSE MERCADO/STANFORD NEWS SERVICE; MCKENZIE: ALEX MCKENZIE; KAHN: LOUIS F. BACHRACH

Bachman had four years earlier received the prestigious Turing Award for his work on a database management system called the Integrated Data Store.

When I interviewed Bachman in 2011, he described the "architectural vision" that he brought to OSI, a vision that was inspired by his work with databases generally and by IBM's Systems Network Architecture in particular. He began by specifying a reference model that divided the various tasks of computer communication into distinct layers. For example, physical media (such as copper cables) fit into layer 1; transport protocols for moving data fit into layer 4; and applications (such as e-mail and file transfer) fit into layer 7. Once a layered architecture was established, specific protocols would then be developed.

OSI vs. TCP/IP



A LAYERED APPROACH: The OSI reference model [left column] divides computer communications into seven distinct layers, from physical media in layer 1 to applications in layer 7. Though less rigid, the TCP/IP approach to networking can also be construed in layers, as shown on the right.

Bachman's design departed from IBM's Systems Network Architecture in a significant way: Where IBM specified a terminal-to-computer architecture, Bachman would connect computers to one another, as peers. That made it extremely attractive to companies like General Motors, a leading proponent of OSI in the 1980s. GM had dozens of plants and hundreds of suppliers, using a mix of largely incompatible hardware and software. Bachman's scheme would allow "interworking" between different types of proprietary computers and networks—so long as they followed OSI's standard protocols.

The layered OSI reference model also provided an important organizational feature: modularity. That is, the layering allowed committees to subdivide the work. Indeed, Bachman's reference model was just a starting point. To become an international standard, each proposal would have to complete a four-step process, starting with a working draft, then a draft proposed international standard, then a draft international standard, and finally an international standard. Building consensus around the OSI reference model and associated standards required an extraordinary number of plenary and committee meetings.

OSI's first plenary meeting lasted three days, from 28 February through 2 March 1978. Dozens of delegates from 10 countries participated, as well as observers from four international organizations. Everyone who attend-

ed had market interests to protect and pet projects to advance. Delegates from the same country often had divergent agendas. Many attendees were veterans of INWG who retained a wary optimism that the future of data networking could be wrested from the hands of IBM and the telecom monopolies, which had clear intentions of dominating this emerging market.

Meanwhile, IBM representatives, led by the company's capable director of standards, Joseph De Blasi, masterfully steered the discussion, keeping OSI's development in line with IBM's own business interests. Computer scientist John Day, who designed protocols for the ARPANET, was a key member of the U.S. delegation. In his 2008 book *Patterns in Network Architecture* (Prentice Hall), Day recalled that IBM representatives expertly intervened in disputes between delegates "fighting over who would get a piece of the pie.... IBM played them like a violin. It was truly magical to watch."

Despite such stalling tactics, Bachman's leadership propelled OSI along the precarious path from vision to reality. Bachman and Hubert Zimmermann (a veteran of Cyclades and INWG) forged an alliance with the telecom engineers in CCITT. But the partnership struggled to overcome the fundamental incompatibility between their respective

worldviews. Zimmermann and his computing colleagues, inspired by Pouzin's datagram design, championed "connectionless" protocols, while the telecom professionals persisted with their virtual circuits. Instead of resolving the dispute, they agreed to include options for both designs within OSI, thus increasing its size and complexity.

This uneasy alliance of computer and telecom engineers published the OSI reference model as an international standard in 1984. Individual OSI standards for transport protocols, electronic mail, electronic directories, network management, and many other functions soon followed. OSI began to accumulate the trappings of inevitability. Leading computer companies such as Digital Equipment Corp., Honeywell, and IBM were by then heavily invested in OSI, as was the European Economic Community and national governments throughout Europe, North America, and Asia.

Even the U.S. government—the main sponsor of the Internet protocols, which were incompatible with OSI—jumped on the OSI bandwagon. The Defense Department officially embraced the conclusions of a 1985 National Research Council recommendation to transition away from TCP/IP and toward OSI. Meanwhile, the Department of Commerce issued a mandate in 1988 that the OSI standard be used in all computers purchased by U.S. government agencies after August 1990.

While such edicts may sound like the work of overreaching bureaucrats, remember that throughout the 1980s, the Internet was still a *research* network: It was growing rapidly, to be sure, but its managers did not allow commercial traffic or for-profit service providers on the government-subsidized backbone until 1992. For businesses and other large entities that wanted to exchange data between different kinds of computers or different types of networks, OSI was the only game in town.

That was not the end of the story, of course. By the late 1980s, frustration with OSI's slow development had reached a boil-



WHAT'S IN A NAME: At a July 1986 meeting in Newport, R.I., representatives from France, Germany, the United Kingdom, and the United States considered how the OSI reference model would handle the crucial functions of naming and addressing on the network.

ing point. At a 1989 meeting in Europe, the OSI advocate Brian Carpenter gave a talk titled “Is OSI Too Late?” It was, he recalled in a recent memoir, “the only time in my life” that he “got a standing ovation in a technical conference.” Two years later, the French networking expert and former INWG member Pouzin, in an essay titled “Ten Years of OSI—Maturity or Infancy?,” summed up the growing uncertainty: “Government and corporate policies never fail to recommend OSI as the solution. But, it is easier and quicker to implement homogenous networks based on proprietary architectures, or else to interconnect heterogeneous systems with TCP-based products.” Even for OSI’s champions, the Internet was looking increasingly attractive.

That sense of doom deepened, progress stalled, and in the mid-1990s, OSI’s beautiful dream finally ended. The effort’s fatal flaw, ironically, grew from its commitment to openness. The formal rules for international standardization gave any interested party the right to participate in the design process, thereby inviting structural tensions, incompatible visions, and disruptive tactics.

OSI’s first chairman, Bachman, had anticipated such problems from the start. In a conference talk in 1978, he worried about OSI’s chances of success: “The organizational problem alone is incredible. The technical problem is bigger than any one previously faced in information systems. And the political problems will challenge the most astute statesmen. Can you imagine trying to get the representatives from ten major and competing computer corporations, and ten telephone companies and PTTs [state-owned telecom monopolies], and the technical experts from ten different nations to come to any agreement within the foreseeable future?”

Despite Bachman’s and others’ best efforts, the burden of organizational overhead never lifted. Hundreds of engineers attended the meetings of OSI’s various committees and working groups, and the bureaucratic procedures used to structure the discussions didn’t allow for the speedy production of standards. Everything was up for debate—even trivial nuances of language, like the difference between “you will comply” and “you should comply,” triggered complaints. More significant rifts continued between OSI’s computer and telecom experts, whose technical and business plans remained at

odds. And so openness and modularity—the key principles for coordinating the project—ended up killing OSI.

Meanwhile, the Internet flourished. With ample funding from the U.S. government, Cerf, Kahn, and their colleagues were shielded from the forces of international politics and economics. ARPA and the Defense Communications Agency accelerated the Internet’s adoption in the early 1980s, when they subsidized researchers to implement Internet protocols in popular operating systems, such as the modification of Unix by the University of California, Berkeley. Then, on 1 January 1983, ARPA stopped supporting the ARPANET host protocol, thus forcing its contractors to adopt TCP/IP if they wanted to stay connected; that date became known as the “birth of the Internet.”

And so, while many users still expected OSI to become the future solution to global network interconnection, growing numbers began using TCP/IP to meet the practical near-term pressures for interoperability.

Engineers who joined the Internet community in the 1980s frequently misconstrued OSI, lampooning it as a misguided monstrosity created by clueless European bureaucrats. Internet engineer Marshall Rose wrote in his 1990 textbook that the “Internet community tries its very best to ignore the OSI community. By and large, OSI technology is *ugly* in comparison to Internet technology.”

Unfortunately, the Internet community’s bias also led it to reject any technical insights from OSI. The classic example was the “palace revolt” of 1992. Though not nearly as formal as the bureaucracy that devised OSI, the Internet had its Internet Activities Board and the Internet Engineering Task Force, responsible for shepherding the development of its standards. Such work went on at a July 1992 meeting in Cambridge, Mass. Several leaders, pressed to revise routing and addressing limitations that had not been anticipated when TCP and IP were designed, recommended that the community consider—if not adopt—some technical protocols developed within OSI. The hundreds of Internet engineers in attendance howled in protest and then sacked their leaders for their heresy.

Although Cerf and Kahn did not design TCP/IP for business use, decades of govern-

ment subsidies for their research eventually created a distinct commercial advantage: Internet protocols could be implemented for free. (To use OSI standards, companies that made and sold networking equipment had to purchase paper copies from the standards group ISO, one copy at a time.) Marc Levilion, an engineer for IBM France, told me in a 2012 interview about the computer industry’s shift away from OSI and toward TCP/IP: “On one side you have something that’s free, available, you just have to load it. And on the other side, you have something which is much more architected, much more complete, much more elaborate, but it is expensive. If you are a director of computation in a company, what do you choose?”

By the mid-1990s, the Internet had become the de facto standard for global computer networking. Cruelly for OSI’s creators, Internet advocates seized the mantle of “openness” and claimed it as their own. Today, they routinely campaign to preserve the “open Internet” from authoritarian governments, regulators, and would-be monopolists.

In light of the success of the nimble Internet, OSI is often portrayed as a cautionary tale of overbureaucratized “anticipatory standardization” in an immature and volatile market. This emphasis on its failings, however, misses OSI’s many successes: It focused attention on cutting-edge technological questions, and it became a source of learning by doing—including some hard knocks—for a generation of network engineers, who went on to create new companies, advise governments, and teach in universities around the world.

Beyond these simplistic declarations of “success” and “failure,” OSI’s history holds important lessons that engineers, policymakers, and Internet users should get to know better. Perhaps the most important lesson is that “openness” is full of contradictions. OSI brought to light the deep incompatibility between idealistic visions of openness and the political and economic realities of the international networking industry. And OSI eventually collapsed because it could not reconcile the divergent desires of all the interested parties. What then does this mean for the continued viability of the open Internet? ■

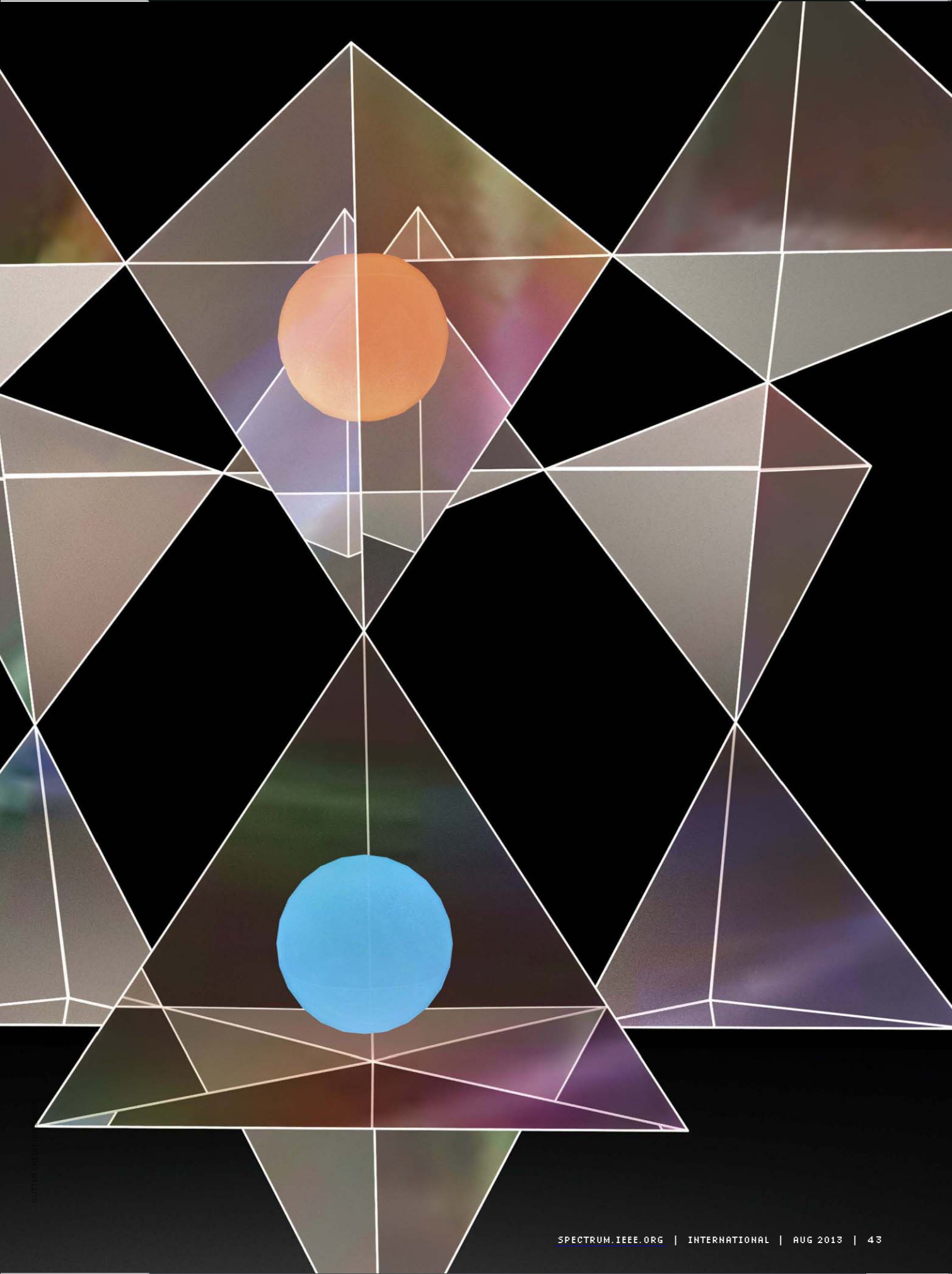
POST YOUR COMMENTS online at
<http://spectrum.ieee.org/osihistory0813>

● THE RACE FOR THE POLE ●

**If magnets don't have to have two poles,
it could lead to an entirely new class of devices**

BY JONATHAN MORRIS ILLUSTRATION BY BRYAN CHRISTIE DESIGN

ONE OF THE VERY FIRST FACTS you learn about electromagnetism—long before you walk into your first physics class—is that every magnet has two poles. Cut a bar magnet in half and you wind up with two magnets, each of which has its own north and south poles. And that's true for every single object in our experience that boasts a magnetic field—whether it's the entire Earth or an iron atom. There are no solitary poles. ¶ Strangely, though, there is no fundamental reason why that has to be the case. In fact, there are a few good reasons to suspect that there might be single-poled magnetic objects—magnetic monopoles—floating about in the universe. If these particles exist, they are probably quite rare, but that hasn't stopped physicists from looking for them. Here's why: If they exist, they could help answer long-standing questions about the nature of the universe, shedding light on the way fundamental forces of nature are tied together. ¶ So over the past few decades, physicists have scoured Antarctic ice and lunar dust and scrutinized rocks culled from ocean beds and polar volcanoes. They've tried to create monopoles in particle colliders and have searched for signatures in cosmic rays that collide with Earth.



In every case so far, the search has come up dry. But a few years ago, my colleagues and I, along with other research groups, came across evidence of something that looks and acts very much like a naturally occurring monopole would. Our monopoles are confined to particular materials, and they arise only when the spins of atoms are aligned in just the right way. But unlike their still elusive, relatively unconstrained brethren, they do offer some hope of new technologies. One day, we might be able to manipulate magnetic charges much as we control the flow of electric charges today. It's almost impossible to predict where this new capability could lead; we could see devices that are capable of performing computations or storing information or energy in entirely new ways. But before we can know what monopoles can do, we must get to the bottom of how they behave.

MOST OF WHAT WE KNOW about magnetic fields dates back to the mid-19th century, when James Clerk Maxwell debuted a set of equations that demonstrated that electric and magnetic forces both stem from a single force: electromagnetism. Maxwell showed there is symmetry in this unification: Changing magnetic fields create electric currents, and moving electric charges create magnetic fields. But this symmetry has its limits. Electric charges can be positive or negative, but magnets are dipolar. They always contain two “magnetic charges”: a north and a south pole.

This asymmetry was more or less just a curiosity until 1931. That's when physicist Paul Dirac showed that the existence of a magnetic version of charge—a magnetic monopole—could help explain a seemingly arbitrary fact: why electrons and other charged particles have only quantized amounts—that is to say, integer multiples—of a fundamental electric charge. That realization offered the possibility that, even if they are far removed from our everyday experience, magnetic monopoles just might exist. More recent theoretical work has revitalized this idea, because the particles also pop up in grand unified theories that attempt to tie fundamental forces of nature together.

These monopoles would be free-floating particles, at liberty to flit about the vacuum of space. But in 2008, three theoretical physicists—Claudio Castelnovo, Roderich Moessner, and Shivaji Sondhi—argued that we should be able to find something like them inside special crystals on Earth.

Such monopoles wouldn't be fundamental particles; they would exist only in the material and would form from the collective behavior of other particles. But they would technically be single magnetic charges, and they would interact with one another in much the same way that fundamental monopoles probably would.

The team proposed looking for these trapped monopoles at temperatures close to absolute zero in spin ice, a peculiar class of materials with ions arranged in four-sided pyramids called tetrahedra. These tetrahedra are stacked together to make a crystal called a pyrochlore.

The atoms at each corner of the pyramids in a pyrochlore are magnetic dipoles. Just like a bar magnet, they have a magnetic field that emerges from one side (what physicists tend to call “north” by convention) and curves around the atom so that it eventually enters the opposite end (“south”). As shorthand, physicists represent this magnetic field as a north-pointing arrow centered on the atom. This arrow is an indicator of magnetism, and it also corresponds to the atom's “spin”—its intrinsic angular momentum. Spin is a quantum mechanical concept: Atoms may not physically spin, but their angular momentum is reflected in other ways, such as the trajectories that daughter particles take when radioactive atoms decay.

Just like bar magnets, these spins interact with one another electromagnetically and try to align themselves so that they are parallel or antiparallel to one another. But what makes spin ices special—and what gives them their name—are the configurations these spins are forced into by the geometry of the crystal.

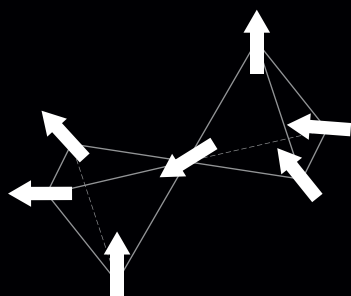
When the temperature of the crystalline material is relatively high, the forces that try to align the spins are easily overwhelmed by thermal fluctuations. The spins

BORN IN PAIRS

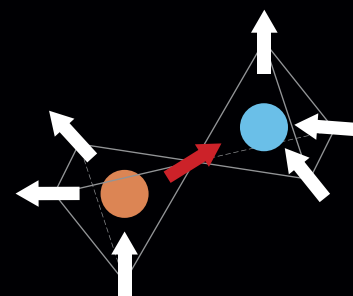
Spin-ice monopoles are created in pairs but can move apart from one another. The aftermath of this migration is one thing that makes them detectable.

THE BIRTH OF A MONOPOLE

©2009 IEEE. ALL RIGHTS RESERVED. PHOTO COURTESY OF CLAUDIO CASTELNOVO, Roderich Moessner, and Shivaji Sondhi.



THE SPINS OF THE FOUR ATOMS in a spin-ice tetrahedron naturally prefer to align in the lowest-energy configuration, with two spins pointing toward the center of the group and two spins pointing away. The two tetrahedra



[left] obey this “ice rule.” If there is enough energy, the pair's shared spin might flip and the ice rule will be violated in both tetrahedra. This creates a pair of monopoles, a “north” [blue] and a “south” [orange].

are oriented at random and can easily change direction. When the material is cooled to just a few degrees above absolute zero, the forces between spins begin to dominate. The spins of the atoms naturally begin to align into the lowest, most stable energy state: a configuration in which two of each tetrahedron's four spins point toward the center of the pyramid and two point outward.

This "two-in, two-out" configuration is called the ice rule, so called because it's analogous to the way that hydrogen bonds form in ordinary water-ice crystals. For each oxygen atom, two hydrogen ions sit close in and two sit farther away. Spin ices tend to obey the ice rule, but the alignment process isn't perfect. Sometimes defects form as the sample cools. In other cases, a bit of thermal energy can cause a spin to flip its orientation. In the end, some tetrahedra wind up with three spins pointing in or out and just one in the opposite direction. Researchers have known about these rule violations for a while, but Castelnovo and his colleagues argued that they might have a much bigger physical significance.

Because magnetic fields extend well beyond each atom, the researchers suggested that we think of the spin on each atom as an extended dumbbell. The connecting bar is centered on the atom's nucleus and ends in two distinct magnetic charges: a north pole at the center of one tetrahedron and a south pole at the center of another. When you look at the center of one of those tetrahedra, the poles from the four corners will either cancel each other out or add together.

In the case where ice rules are obeyed, the two north poles and two south poles cancel each other out. But here's where it gets interesting: When the ice rules are not obeyed—if, for example, there are three spins pointing inward and one pointing outward—then the

three north poles and one south pole in the center will give rise to a single, north magnetic pole. Voilà! There's your monopole. Similarly, a one-in, three-out configuration would make a single, south magnetic monopole.

And that's not all. Because each vertex of a tetrahedron is also shared with a neighbor, the spin that violates the ice rule in one tetrahedron would also do so in another, creating a second monopole one tetrahedron away with the opposite polarity. So, for every monopole, there will be an oppositely charged monopole. If that same spin flips, or rotates 180 degrees, the spin-ice rules will be fulfilled and the pair will essentially be annihilated.

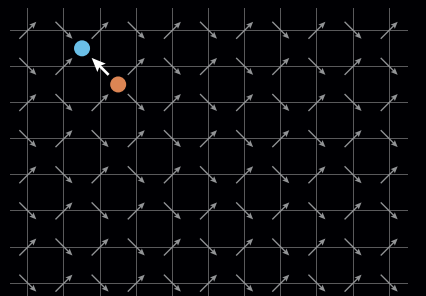
To be considered true magnetic particles and not just magnetic anomalies, these spin-ice monopoles should be able to move. And indeed, Castelnovo and his colleagues noted that two monopoles should be able to separate from one another and move independently through the spin ice. Flip another spin in a monopole's tetrahedron and the spins will once again balance out there but become unbalanced in a neighbor. The monopole will have effectively jumped from one tetrahedron to another. That raised the tantalizing possibility of creating the magnetic analogue of electric current: "magnetic current."

ONCE THIS PROPOSAL APPEARED in *Nature* in January of 2008, it didn't take long for experimentalists around the world to go on the hunt. At the time, I was working as a postdoc at the Helmholtz Center Berlin for Materials and Energy. My collaborators and colleagues—Alan Tennant and Santiago Grigera—had already been working on spin ices for several months. They had been examining some anomalous behavior that set in when a particular spin ice—a compound made of oxygen, titanium, and the rare earth element dysprosium—was cooled to below 0.6 kelvin. When the theory paper came out, they changed focus, and I joined my colleagues—Tennant, Grigera, and a third team member, Kirrily Rule—to hunt for evidence of magnetic monopoles.

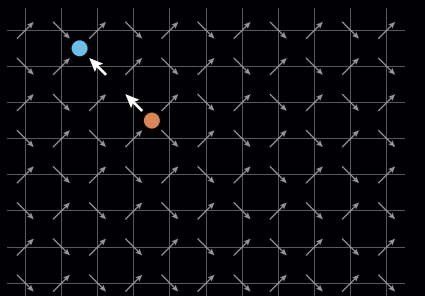
We used a powerful probe we already had on hand: neutron beams. Because they have no electric charge, neutrons pass easily through most of the space in a material, scattering only when they collide with the tiny nuclei at the center of atoms. But they also have intrinsic magnetic moments, so they can be deflected by the magnetic fields of atoms. Shine a beam of neutrons through a material and the pattern they make on the other side

MONOPOLE MIGRATION

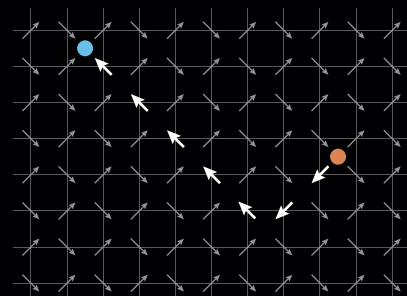
PHOTOGRAPH BY JAMES HARRIS FOR SPECTRUM



MONOPOLES CAN BE MOVED with the application of an external magnetic field. This 2-D representation of the spin ice illustrates the basic configurations. Two monopoles begin at the centers of adjoining intersections of four spins.



A MAGNETIC FIELD can encourage adjacent spins to flip, restoring the spin-ice rules to the original tetrahedron and advancing the poles in opposite directions.



OVER TIME, two monopoles can separate from each other by a considerable distance. They leave a trail, or tube, of flipped spins between them, a structure that can be detected by shining neutrons through the sample.

can tell you certain things about the magnetic structure that caused the neutrons to diffract. It's the standard way physicists study subjects as wide-ranging as the structure of proteins, strains in jet engine blades, and the magnetic interactions in superconductors.

In a spin ice, neutrons are particularly well suited to picking up on correlations between spins—how closely related the alignment of one spin is to another. We think of these alignments as “polarization tubes”—chains of spins that by convention flow from south to north. In an ordinary tetrahedron that obeys the ice rules, you'll have two such tubes flowing through the center; each one enters the tetrahedron by a spin with its north side pointing in and exits by one with its north side pointing out. When the ice rules are broken, just one polarization tube will flow all the way through the tetrahedron and two other polarization tubes will terminate at the center. Those endpoints are where we should expect to see monopoles. Otherwise we might simply see a material containing polarization tubes without endpoints—essentially just loops.

When neutrons scatter off polarization tubes, they leave distinctive patterns in the locations of neutrons that hit our detectors. But inferring the structure responsible for this pattern is difficult to do. If you could visualize all the polarization tubes in a spin ice, you'd have what would look like a bowl of spaghetti. If you tried to follow a single strand through a pile, you'd quickly become lost in a 3-D knot.

To simplify matters, our group devised a way to, in effect, pull on that tangle of spaghetti with a fork, separating out the polarization tubes by applying a magnetic field. When a strong enough magnetic field is applied to a spin ice, all of the spins orient themselves along the field, and because they are all pointing in the same direction, the polarization tubes are erased. As the magnetic field is lowered back down, spins begin flipping against the magnetic field as the strength of their interactions starts to overcome the external magnetic field. These flipped spins form short “strings”—sections of polarization tube that are easier to see in the neutron data. When we looked at the data as a function of the applied magnetic field, we found that the strings grow along the field, but they also meander a bit.

This is consistent with what we would expect to see if monopoles were present and had enough energy to migrate away from one another, leaving behind

a string of flipped spins. But we needed a way to confirm there were in fact monopoles at the end of the strings.

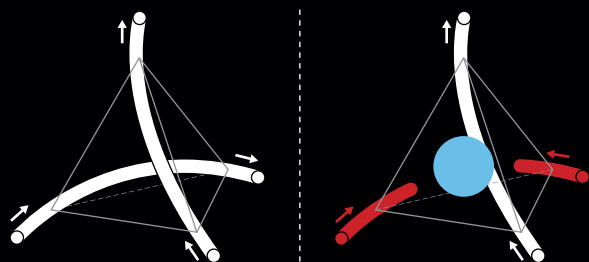
Fortunately, some colleagues of ours in Berlin—Clemens Czternasty, Bastian Klemke, and Michael Meissner—were working on just such an experiment. They were investigating the spin ice's heat capacity, a measure of how much energy is needed to change a material's temperature. Heat capacity is linked to the number of different configurations a material can take on. For water, it corresponds to the number of possible arrangements of hydrogen and oxygen atoms. For a spin ice, it's related to the number of spin orientations. The number of possible configurations is linked to how much energy is available to the materials, so heat capacity is a very sensitive measure of the state of the system.

The Berlin team found that very little heat was needed to raise the temperature at very low sample temperatures. The drop-off in heat capacity as the temperature goes down was so sharp that the team initially thought something was wrong with the equipment. But repeated measurements confirmed the result. When we presented the data to Castelnovo and his colleagues, they first attempted to make sense of it with a simple model of the material that was based on spins alone. But they found that the only way they could get a good fit to the data was to posit the existence of monopoles that have their own physical interactions. It was estimated that we created about 600 monopoles per cubic centimeter.

Armed with this evidence, our team started preparing a research paper. But we soon found out that other teams were also on the case. In March 2009, a pair of theoretical physicists based in France reported in *Nature Physics* that they could fit some old magnetic data better if the existence of magnetic monopoles was assumed.

SPIN STRUCTURE

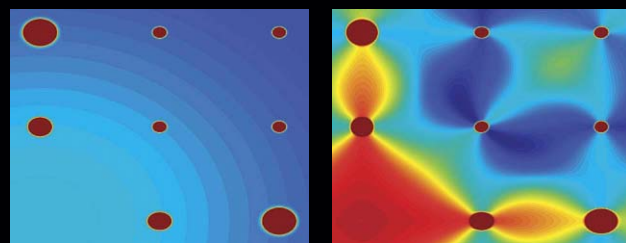
PHYSICISTS CAN VISUALIZE spin-ice structure through polarization tubes.



ONE WAY THAT PHYSICISTS VISUALIZE spin-ice structure is through polarization tubes. These connect spins from south to north, much like a conga line. A spin-ice tetrahedron that obeys the ice rule will have two polarization tubes running through it [left]. Tubes terminate when a monopole is present [right].

BENDING NEUTRONS

PHYSICISTS CAN LOOK for polarization tubes, or strings, that link monopoles by shining neutrons through a spin-ice sample.



PHYSICISTS CAN LOOK for polarization tubes, or strings, that link monopoles by shining neutrons through a spin-ice sample. These idealizations show what the resulting neutron-scattering data might look like if a sufficiently strong magnetic field is used to align all the spins so that no strings exist [left] and if this magnetic field is turned off so strings can form [right].

This made a bit of a media splash, but we realized that we had an even more detailed story to tell, so we remained undeterred. Then, in May 2009, I saw a presentation on a hunt by Hiroaki Kadowaki's group at the Tokyo Metropolitan University, at the International Conference on Neutron Scattering. By the time we submitted our paper to *Science* in July, we found out that two other papers were also under consideration by the magazine.

In the end, three experimental teams all reported evidence of magnetic monopoles in quick succession. Our *Science* paper appeared in September alongside one from a British-French group led by Tom Fennell, then at the Institut Laue-Langevin. That team reported evidence that strings lengthen in spin ice as the temperature drops, which suggested that if monopoles were present, they were migrating away from one another in the sample. Around the same time, Kadowaki's team reported neutron scattering results that seemed to match simulations of scattering as a function of monopole density.

ONE OF THE QUESTIONS I was asked a lot after our *Science* paper came out was whether the monopoles we'd found in spin ice could be considered "real" magnetic monopoles.

They're certainly not the ones that particle physicists have been dreaming about for decades. The magnetic monopoles in spin ice can't exist in free space. In a sense, they're not real particles at all. Physicists prefer to call them quasiparticles: particle-like entities that are artifacts of the collective behavior of other particles. Quasiparticles are common. Many modern semiconductor devices, for example, rely on the flow of holes—positively charged "particles" that are really just the absence of electrons.

If these monopoles are not the objects searched for by particle physicists, can we justifiably call them magnetic monopoles? The heat capacity measurements our team performed suggest that spin-ice monopoles really do seem to carry a magnetic charge, and their interactions mimic the basic way that positive and negative magnetic charges would interact in vacuum.

At the same time, they're trapped particles. And that's raised a number of interesting problems. In 2011, one of my Berlin colleagues, Klemke, realized that all the published heat capacity data below 0.6 K were different. What's more, the magnetization measurements in that temperature range don't agree with theory. Some theorists have suggested that these inconsistencies exist because, at low temperatures, there is little energy available to flip spins. As it equilibrates, a spin ice could get trapped in an especially unfavorable configuration that is difficult to escape from without some very well-coordinated, simultaneous spin flipping. That means it will take a long time for the material to find the most energetically favorable configuration. In fact, in April, a Canadian group led by David Pomaranski reported that

at very low temperatures, spin ice can take weeks to relax; at just one degree above absolute zero, it might take a hundredth of a second.

Because we're still sorting out the basic physics, we've only just begun to start thinking about applications. One intriguing possibility is "magnetricity"—a magnetic form of electricity. With the right magnetic field, it might be possible to draw monopoles through spin ice in much the same way that voltage is used to draw current through wires. The peculiarities of the spin-ice geometry won't make magnetic current nearly as simple, however. Because moving monopoles leave a trail of flipped spins behind them, it's impossible to drive a second monopole along the very same path. That rules out a monopole direct current, but it does leave the window open for monopole alternating current and devices.

Because the spins can flip in a spin ice, the material might also be able to screen a magnetic field and store magnetic energy much as a dielectric does. What's more, there is some evidence to suggest that it might be possible to dope the material, similarly to the way that impurities are added to a semiconductor to boost its charge-carrying properties. Last year, experiments showed that monopoles could be made to move more slowly through a spin ice if more magnetic ions were added to the mix. If the spin ice can be doped to alter the speed of monopoles through the material, it's not unreasonable to imagine magnetic analogues to basic electrical device components like capacitors and junctions.

But there are a few big stumbling blocks we'd have to overcome before we could get anywhere close to making practical magnetronic devices. One is the purity of the samples: Even tiny structural defects can block the flow of monopoles. The other is temperature. To freeze a spin ice and create monopoles, you must cool such materials to very low temperatures, typically on the order of one degree above absolute zero—about 1 K, or -272 °C. It's unclear whether we can create a material that can undergo a spin-ice transition at more practical temperatures. The magnetic moments of its atoms, or the interactions between them, would have to be enormous to counteract the scrambling effects of thermal energy at higher temperatures, and no such material is currently known to exist.

One alternative monopole system that has emerged is "artificial spin ice." These man-made materials are two-dimensional systems that can be made from nanoscale patches—either islands or wires—of a ferromagnetic material such as cobalt. A patch's dimensions are chosen so the magnetic moments of the atoms inside it point toward a vertex between neighbors. If these patches are arranged in either a honeycomb or a square structure, you can get their magnetic moments to obey ice rules. A square lattice will obey the same two-in, two-out rules as ordinary spin ice. In a honeycomb lattice, where three patches meet at each vertex, there is always an excess magnetic charge.

How do we know that these charges exist? Incredibly, they can be imaged using magnetic force microscopy. In 2010, an Imperial College London team observed artificial magnetic monopoles and was even able to watch them move under influence of a magnetic field. These materials have the benefit of being stable at room temperature, and some researchers have proposed that they might make useful memory devices. But artificial spin-ice patches currently have dimensions in the 100-nanometer range, making them gigantic by existing industry standards.

Still, it's early days for the magnetic monopole, and I wouldn't rule out powerful applications that have yet to be imagined. A magnetronic revolution might be just over the horizon. ■

POST YOUR COMMENTS online at <http://spectrum.ieee.org/monopoles0813>

TOSHIBA
Leading Innovation >>>

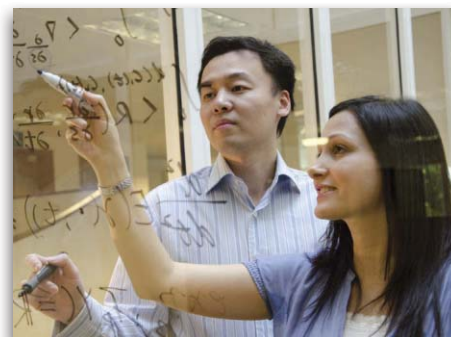
Toshiba Research Europe Limited

Senior/Principal Research Engineer**Bristol, UK**

Toshiba's Telecommunications Research Laboratory has an established track-record of innovation in wireless communications. Situated in one of England's most attractive city centres, it works closely with numerous partners, including nearby Bristol University; it files patents, publishes in leading peer-reviewed journals and conferences, and transfers advanced concepts and prototypes to Toshiba's business units.

The Lab's *Physical Layer Research Group* is looking for experienced applicants to engage in fundamental research as well as the design of communication technologies for reliable, secure, energy-efficient networks.

Successful applicants will have a PhD, some years of post-doctoral experience in a high-quality research environment and a track-record of excellence, likely evidenced through publication in leading journals and conferences. Typical experience might include: PHY layer modulation (e.g. OFDM, SC-FDMA), wireless sensor & *ad hoc* networks, multi-antenna techniques (e.g. MU-MIMO, massive MIMO, space-time processing), physical layer security, cooperative communication & relays, large-scale networks, communication theory, optimization theory.



A permanent position is offered (subject to probation). Salary £40,000-£50,000 depending on experience. Closing date 16th September 2013.

Further information or applications:
<http://www.toshiba.eu/TRL-Jobs>
recruitment@toshiba-trel.com

**THE HONG KONG UNIVERSITY OF SCIENCE AND TECHNOLOGY****Head of the Department of
Computer Science and Engineering**

The Hong Kong University of Science and Technology (HKUST), opened in October 1991, comprises four Schools: Science, Engineering, Business & Management, and Humanities & Social Science. The University's mission is to advance learning and scholarship; to promote research, development, and entrepreneurship, and to contribute to the region's economic and social development.

The School of Engineering, the largest School of the University, currently enrolls about 38% of the University's total undergraduate and postgraduate students of approximately 12,600. It comprises six departments: Chemical & Biomolecular Engineering, Civil & Environmental Engineering, Computer Science & Engineering, Electronic & Computer Engineering, Industrial Engineering & Logistics Management, and Mechanical Engineering.

The Department of Computer Science & Engineering (CSE) currently has 44 faculty members, teaching about 560 undergraduate students and 180 postgraduate research students. The Department conducts comprehensive teaching and research programs in both basic and applied aspects of Computer Science & Engineering. The academic degrees offered by the Department are: BEng, MSc, MPhil and PhD. Research activities in the Department are broadly categorized into artificial intelligence; data, knowledge and information management; networking and computer systems; software technologies; theoretical computer science; vision and graphics. For more information, please visit the University and Department websites available on <http://www.ust.hk/> and <http://www.cse.ust.hk/> respectively.

Applications/nominations are invited from well-qualified and accomplished scholars for the position. In addition to extensive teaching and research experience, the successful candidate must have demonstrated leadership qualities necessary to lead and manage the Department in its diverse academic and administrative functions and to interact effectively with industry and commerce.

Salary will be highly competitive with generous benefits. Applications/nominations together with detailed curriculum vitae and the names and addresses/fax numbers/email addresses of three referees should be sent to Professor Chung Yee Lee, Chair of Search Committee for Headship of CSE, c/o School of Engineering, HKUST, Clearwater Bay, Kowloon, Hong Kong [Fax No.: (852) 2358 1458, e-mail: dhcse@ust.hk] before **Tuesday, 1 October 2013**.

HKUST is committed to increasing the diversity of its faculty and has a range of family-friendly policies in place.

(Information provided by applicants will be used for recruitment and other employment-related purposes.)

**FACULTY
SEARCH****ShanghaiTech University****School of Information Science and Technology**

The School of Information Science and Technology (SIST) in the new ShanghaiTech University invites applications to fill multiple tenure-track and tenured positions. Candidates should have an exceptional academic record or strong potential in frontier areas of information sciences.

ShanghaiTech is founded as a world-class research university for training future scientists, entrepreneurs, and technology leaders. Besides keeping a strong research profile, successful candidates should also contribute to undergraduate and graduate education within SIST.

Compensation and Benefits:

Salary and startup fund are highly competitive, commensurate with academic experience and accomplishment. ShanghaiTech also offers a comprehensive benefit package which includes housing. All regular faculty members are hired within ShanghaiTech's new tenure-track system commensurate with international practice and standards.

Academic Disciplines:

We seek candidates in all cutting edge areas of information science and technology. Our recruitment focus includes, but is not limited to: computer architecture and technologies, nano-scale electronics, high speed and RF circuits, intelligent and integrated signal processing systems, computational foundations, big data, data mining, visualization, computer vision, bio-computing, smart energy/power devices and systems, next-generation networking, as well as inter-disciplinary areas involving information science and technology.

Qualifications:

- Well developed research plans and demonstrated record/potentials;
- Ph.D. (Electrical Engineering, Computer Engineering, Computer Science, or related field);
- A minimum relevant research experience of 4 years.

Applications:

Submit (all in English) a cover letter, a 2-page research plan, a CV including copies of 3 most significant publications, and names of three referees to: sist@shanghaitech.edu.cn.

Deadline: September 15th, 2013 (or until positions are filled). We have 10 positions for this round of faculty recruitment.

For more information, visit <http://www.shanghaitech.edu.cn>.



Joint Institute of Engineering



FACULTY POSITIONS AVAILABLE IN ELECTRICAL/COMPUTER ENGINEERING

Sun Yat-sen University & Carnegie Mellon University are partnering to establish the **SYSU-CMU Joint Institute of Engineering (JIE)** to innovate engineering education in China and the world. The mission of the JIE is to nurture a passionate and collaborative global community and network of students, faculty and professionals working toward pushing the field of engineering forward through education and research in China and in the world.

JIE is seeking **full-time faculty** in all areas of electrical and computer engineering (ECE). Candidates should possess a doctoral degree in ECE or related disciplines, with a demonstrated record and potential for research, teaching and leadership. The position includes an initial year on the Pittsburgh campus of Carnegie Mellon University to establish educational and research collaborations before locating to Guangzhou, China.

This is a worldwide search open to qualified candidates of all nationalities, with an internationally competitive compensation package for all qualified candidates.

PLEASE VISIT: sysucmuje.cmu.edu for details



SHUNDE INTERNATIONAL

Joint Research Institute



RESEARCH STAFF POSITIONS AVAILABLE IN ELECTRICAL/COMPUTER ENGINEERING

SYSU-CMU Shunde International Joint Research Institute (JRI) is located in Shunde, Guangdong. Supported by the provincial government and industry, the JRI aims to bring in and form high-level teams of innovation, research and development, transfer research outcomes into products, develop advanced technology, promote industrial development and facilitate China's transition from labor intensive industries to technology intensive and creative industries.

The JRI is seeking **full-time research faculty** and **research staff** that have an interest in the industrialization of science research, which targets electrical and computer engineering or related areas.

Candidates with industrial experiences are preferred.

Applications should include a full CV, three to five professional references, a statement of research and teaching interests, and copies of up to five research papers.

Please submit the letters of reference and all above materials to the address below.

Application review will continue until the position is filled.

EMAIL APPLICATIONS OR QUESTIONS TO: sdjri@mail.sysu.edu.cn

SUN YAT-SEN UNIVERSITY

Carnegie Mellon University



國立交通大學
National Chiao Tung University

College of Photonics, National Chiao Tung University

invites applications for faculty positions at the Assistant / Associate / Full Professor ranks. The starting date of faculty positions is Feb. 1, 2014 or Aug. 1, 2014. The applicants with demonstrated record in research, commitment to teaching in graduate level, and ability to teach in English, can choose one of the following Institutes:

Institute of Photonic System
Institute of Lighting and Energy Photonics
Institute of Imaging and Biomedical Photonics

National Chiao Tung University has long-lasting been a leading research university in Taiwan. For more information, please see the website of <http://www.nctu.edu.tw/> and <http://www.cop.nctu.edu.tw/>.

Interested applicants should send a cover letter, curriculum vitae, copies of representative publications, photocopy of diploma, research summary and teaching portfolio, and an application form (<http://www.cop.nctu.edu.tw>) with 3 letters of recommendation sent to "Faculty Recruitment Committee, College of Photonics, National Chiao Tung University (Tainan Campus), No.301, Gaofa 3rd Road, Guiren District, Tainan City 711, Taiwan (R.O.C.)" before Nov. 30, 2013

Contact: Ms. Alva Hsu (alvahsu@nctu.edu.tw)

Power Electronics vacancies in the new **National Center for Power Electronics and Energy**: Sun Yat-Sen University, Guangzhou, China, a top university situated in a dynamic region aiming at technology innovation announces openings of Associate and Assistant Professors, and Post-docs. The laboratory covers 300sqm; with an initial budget of 8,000,000, updated equipment was purchased. It aims to be a large world center in the field, performing cutting-edge research in energy conversion, allowing the staff to get international visibility and reputation. Candidates with background in all fields of power electronics are looked for. Submit CV to a.ioinovici@gmail.com

IEEE Open Access

Unrestricted access to today's
groundbreaking research

Learn more about IEEE Open Access:
www.ieee.org/open-access



Innovation doesn't just happen.
Read first-person accounts of
IEEE members who were there.

IEEE Global History Network
www.ieeeahn.org



Technology insight on demand on IEEE.tv

Internet television gets a mobile makeover

A mobile version of IEEE.tv is now available, plus a new app can also be found in your app store. Bring an entire network of technology insight with you:

- Generations of industry leaders.
- The newest innovations.
- Trends shaping our future.

Access award-winning programs about the what, who, and how of technology today.

Go mobile or get the app.
www.ieee.tv



WE'RE NOT JUST TALK – WE'RE THE CONVERSATION.

IEEE Spectrum covers tech news that keeps engineers talking.

Our award-winning newsletters and blogs keep the conversation going with up-to-the minute coverage, all written and reviewed by the world's foremost technology experts!



tech alert

Ground-breaking technology and science news.

robotics news

Advances and news in robotics, automation, control systems, interviews with leading roboticists, and more.

energywise

News and opinions from industry experts on power and energy, green tech and conservation.

computerwise

News and analysis on Software, Systems and IT.

test+ measurement

News about T&M industry, products and processes.

Keep up with the conversation. Subscribe today at spectrum.ieee.org/newsletters.

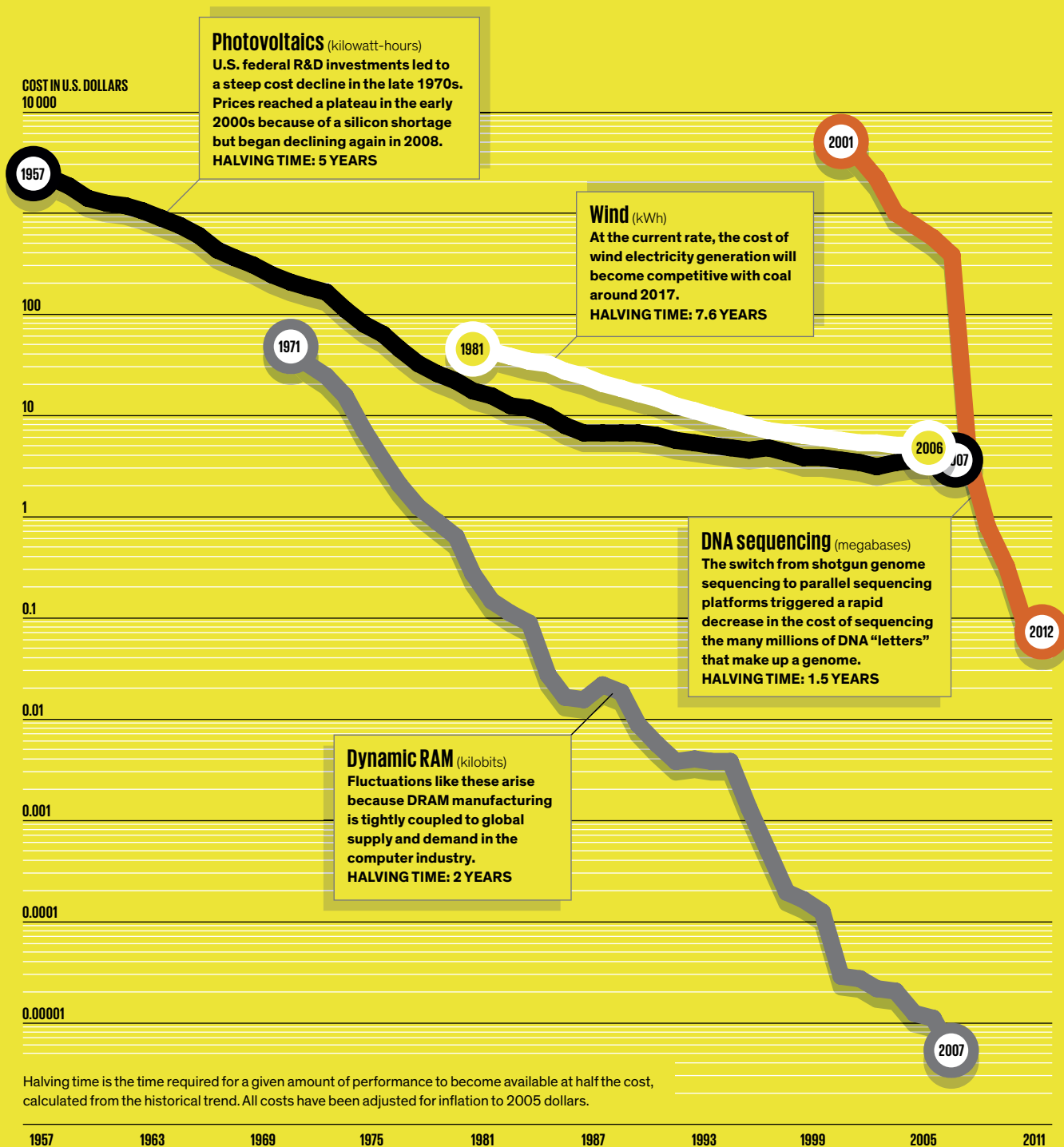
IEEE
SPECTRUM

DATAFLOW_

TECH TRAJECTORIES

EXPONENTIAL IMPROVEMENT ISN'T LIMITED TO THE SEMICONDUCTOR INDUSTRY

We're all familiar with **Moore's Law**, which takes an inexorable view of technological progress, with the number of components on an integrated circuit doubling like clockwork every 18 months or so. But do other technologies follow a similar pattern of exponential improvement? • To find out, researchers at the Santa Fe Institute and MIT looked at 62 technologies across a broad range of industries, merging price and performance data from multiple sources. They found that Moore's Law—like doubling serves as a fair predictor of progress, but not without hiccups. —EMILY ELERT

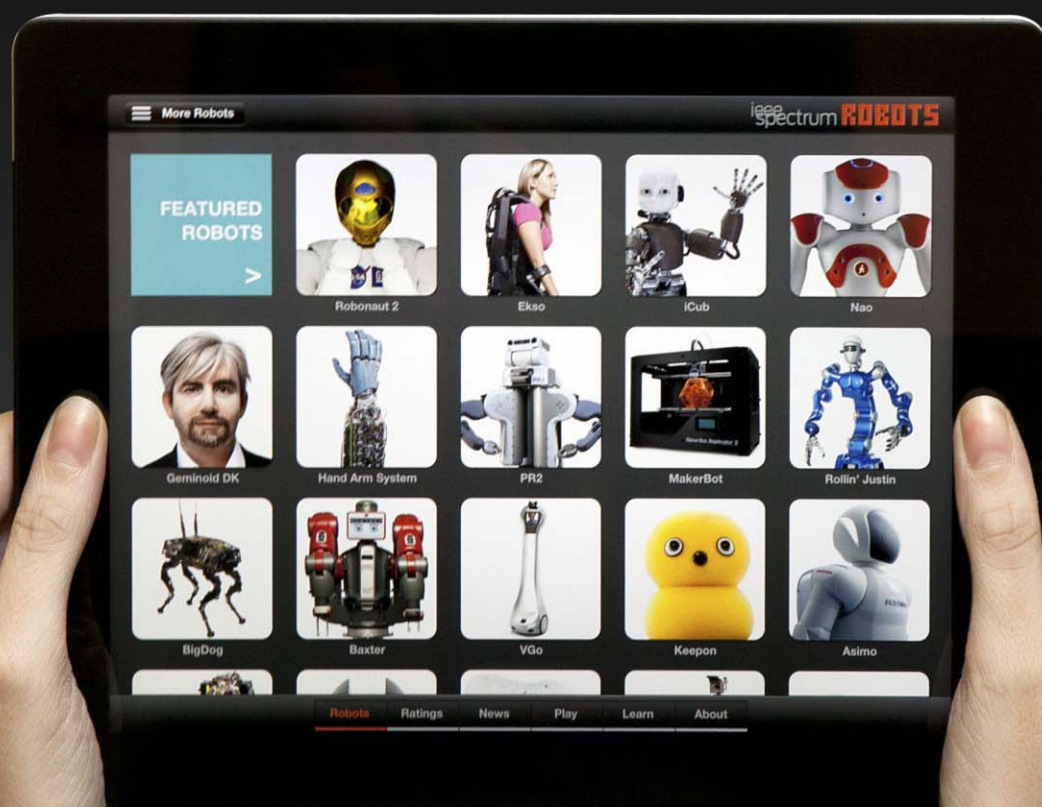
SOURCE: PERFORMANCE CURVE DATABASE ([HTTP://PCDB.SANTAFE.EDU](http://PCDB.SANTAFE.EDU))

“Delightful” - Wired “Robot heaven” - Mashable

Welcome to the world of

ROBOTS

For iPad



Get the app now:
robotsforipad.com



Download on the
App Store



MathWorks®

Accelerating the pace of engineering and science

क्या आप MATLAB बोलते हैं?

Over one million people around the world speak MATLAB.

Engineers and scientists in every field from aerospace and semiconductors to biotech, financial services, and earth and ocean sciences use it to express their ideas.

Do you speak MATLAB?



*Saturn's northern latitudes
and the moon Mimas.
Image from the
Cassini-Huygens mission.*

*Related article at
mathworks.com/ltc*



MATLAB®

The language of technical computing

PHOTO: European Space Agency ©2010 The MathWorks, Inc.