

zation per unit volume or the depth of material that is magnetized—only the product of those two parameters. When mission scientists make a reasonable—though admittedly arbitrary—estimate of the magnetization, such as 20 amperes per meter (within the range of the most intensely magnetized rocks in Earth's crust), they find that the crust must be magnetized to a depth of 30 km, far exceeding the 1-km thickness of Earth's crust. Ron Merrill of the University of Washington points out that the Martian field may well have changed or even reversed during the cooling of such a thick layer, imparting different magnetizations throughout; he believes that, if plate tectonics did operate on Mars, it had quite a different character from what we see on Earth.

Possible explanations

The magnetic stripes seen on Mars are not as definitive as those extending out from the midocean ridges on Earth. Of course, terrestrial researchers can make measurements more easily than can those experimenters trying to gather data on our planetary neighbor. On Earth, researchers have found that the width of a seafloor region with a given magnetic orientation is on the order of tens of kilometers. On Mars, the bands were found to be 200 km wide—the limit of the spatial resolution from an instrument that can get no closer to the surface than 100 km. The measurements can't resolve any structure finer than this. Also, on Earth researchers have been able to date the seafloor around Earth's midocean ridges, thereby confirming the plate tectonics predictions that the crust is older farther from the ridge. If

one could similarly determine the dates of the Martian crust, it would certainly help to test the conjecture about plate tectonics, but such measurements seem out of the question.

What else besides plate tectonics could explain the extensive magnetic pattern that has been seen on Mars? One possibility described by Acuña is that the Martian surface began as a thin, uniformly magnetized crust that was cracked by the same internal rises that gave birth to some of the planet's large volcanoes. With the pressure from below, the surface may have cracked, much like "a muffin in the oven." (Such aligned cracks, called *graves*, have been seen on Mars's surface.) The crustal bar magnets would have broken in half, giving rise to bipolar fields. Later processes may have spread the fissures apart and filled in the gaps. With this scenario, however, it's hard to explain the enormous extent of the observed features. Another possible explanation for the magnetic bands is that the crust was uniformly magnetized at some point and acquired apparent flips in the field direction as a result of crustal folding, but then one must ask, What could cause folding on such a large and regular scale? A third explanation is that the magnetization stems from chemical rather than physical processes.

The hypothesis of plate tectonics on Mars is not entirely new. Back in 1994, before the MGS flew, Norman Sleep of Stanford University proposed that plate tectonics in the northern hemisphere might account for the formation of the lowlands there.⁴ He speculated that a subduction zone might once have paralleled the line of three large volcanoes there.

Serendipity

The very existence of the recent data that revealed intriguing patterns of magnetization on Mars is an unexpected boon: They were taken during a so-called aerobraking maneuver designed to use the friction of Mars's atmosphere to slow the MGS and change its orbit from highly elliptical to circular. While the MGS was still in its elliptical orbit, it passed about 100 km above Mars's surface once in each circuit. By contrast, in its circular orbit, the spacecraft orbits 400 km above the planet's surface. While at the lower altitude, the MGS instruments picked up the magnetic field intensity with 16 times the sensitivity and four times better spatial resolution than is possible in the 400-km orbit. The aerobraking lasted longer than expected so that researchers accumulated ten times more data in this mode than they had counted on. Unfortunately, the aerobraking phase did not allow complete surface coverage nor measurements at a constant altitude, so the measurements at 400 km, now under way, will be a useful complement.

BARBARA GOSS LEVI

References

1. M. H. Acuña, J. E. P. Connerney, N. F. Ness, R. P. Lin, D. L. Mitchell, C. W. Carlson, J. McFadden, K. A. Anderson, H. Rème, C. Mazelle, D. Vignes, P. Wasilewski, P. Cloutier, *Science* **284**, 790 (1999).
2. J. E. P. Connerney, M. H. Acuña, P. Wasilewski, N. F. Ness, H. Rème, C. Mazelle, D. Vignes, R. P. Lin, D. L. Mitchell, P. Cloutier, *Science* **284**, 794 (1999).
3. M. H. Acuña *et al.*, *Science* **279**, 1676 (1998).
4. N. H. Sleep, *J. Geophys. Res.* **99**, 5699 (1994).

Traveling-Wave Thermoacoustic Heat Engines Attain High Efficiency

With continued concerns about greenhouse gases and ozone-depleting refrigerants, the quest for efficient and environmentally benign engines, heat pumps, and refrigerators remains important. For 20 years, one field of exploration has been thermoacoustics, which involves the interplay between thermodynamics and sound waves. Thermoacoustic engine research has focused almost exclusively on standing-wave engines. Now Scott Backhaus and Greg Swift of Los Alamos National Laboratory have reported a new implementation of a *traveling-wave* thermoacoustic engine—a pistonless Stirling engine that's almost twice as efficient as standing-wave

 A new design efficiently converts heat into useful work without any moving parts.

thermoacoustic engines.¹

In this new engine, traveling sound waves in a gas pass through a tightly packed porous medium called a regenerator. One end of the regenerator is kept at ambient temperature, while the other end is heated, thereby establishing a temperature gradient along the regenerator. Gas molecules are kept at the local regenerator temperature as they oscillate back and forth. As a result, parcels of gas expand as they move toward the hotter end of the

regenerator and contract as they move toward the colder end. This cyclic heat transfer, combined with the alternating compression and expansion of the gas produced by the acoustic waves, produces a net increase in the acoustic energy of the sound wave.

A question of phasing

The pressure and velocity of acoustic waves in a gas have rough analogies in AC electric circuits: The pressure resembles the voltage, and the velocity the current. To get any work out of an AC circuit, the voltage and current must not be 90° out of phase. But in a standing-wave thermoacoustic engine (described by Swift in *PHYSICS*

TODAY, July 1995, page 22), the phase between the pressure and velocity is very close to 90° . Consequently, such an engine must rely on an irreversible process for transferring heat to the gas and converting it into work, which limits the engine's efficiency: Current standing-wave thermoacoustic engines typically achieve no better than about 20% of the maximum Carnot efficiency.

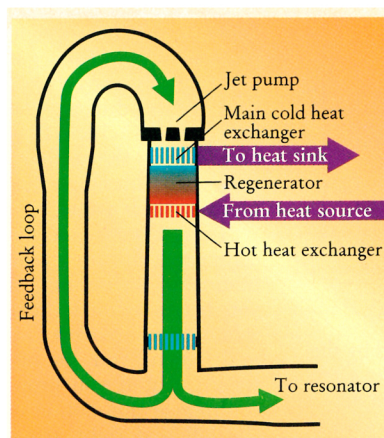
Traveling acoustic waves, in contrast, have their pressure and velocity in phase with each other. Peter Ceperley of George Mason University noted 20 years ago that when traveling waves pass through a regenerator, the thermodynamic cycle of compression, heating, expansion, and cooling that the gas undergoes is the same as in a Stirling engine, where mechanical pistons establish the proper phasing of the gas motion.² With the gas velocity and pressure in phase, a traveling-wave acoustic engine can use a reversible, much more efficient heat transfer process. It therefore does not suffer from the inherent efficiency limitations of standing-wave engines.

Viscous dissipation and other losses have plagued the experimental implementation of traveling wave engines, and the high expectations for these engines are only now beginning to be realized. Just last year, a team led by Taichi Yazaki of the Aichi University of Education in Aichi, Japan, first demonstrated that such an engine can be self-sustaining.³ Now Backhaus and Swift have gone the next step, showing that traveling-wave acoustic engines can have high efficiencies.

Engine design

Five meters long and welded together from steel pipe, the Los Alamos engine consists of a looped tube attached to the small end of a baseball-bat-shaped acoustic resonator. The loop, illustrated in the accompanying figure, contains the heart of the engine. Within the loop, the regenerator—which looks like a tightly packed cylindrical pile of window screen mesh—has its upper end kept at ambient temperature by a water-cooled heat exchanger; heat applied to the lower end provides the energy for the engine. The resonator establishes the operating frequency of 80 Hz.

The loop itself acts like the feedback network in an electronic amplifier circuit, explains Backhaus. "By configuring the acoustic impedances in the loop for positive feedback with the proper traveling-wave phasing at the regenerator, the engine becomes an amplifier for traveling waves." When the apparatus is filled with 30 atmospheres of helium and heat is applied, traveling acoustic waves are generated with peak-to-peak amplitudes of up to 6 atm.



A THERMOACOUSTIC STIRLING ENGINE developed at Los Alamos uses heat to amplify traveling acoustic waves. Sandwiched between cold and hot heat exchangers, a regenerator transfers energy to sound waves as they propagate through the engine's helium gas. Part of the acoustic power produced in the regenerator is fed back around the closed loop; the rest goes to a resonator where it can be tapped to perform work. The jet pump eliminates any net mass circulation around the loop. (Adapted from ref. 1.)

The prototype engine transformed up to 30% of the input heat into acoustic power delivered to the resonator—up to 700 W. This efficiency, which corresponds to over 40% of the Carnot efficiency, already rivals internal combustion engines, and further increases are in store.

Reducing losses

The researchers incorporated several key improvements into their engine to achieve their efficiency results. The first step was to keep the feedback loop containing the regenerator much shorter than one-quarter of an acoustic wavelength. Making the loop short decreases the surface area, thereby reducing the losses around the loop due to viscous drag and thermal hysteresis at the walls.

Furthermore, with a short feedback

loop, the acoustic impedance that the traveling wave experiences is dominated by the loop instead of the regenerator. Backhaus and Swift were therefore able to control the impedance, allowing them to maintain the desired phase relationship between the pressure and velocity as well as to get high pressure oscillations while keeping the velocity amplitude small, which further reduced viscous losses.

An unanticipated problem was that there was a large nonzero average mass flow circulating around the loop, which had a dramatic effect on their engine's performance: It produced an additional heat load that cut the engine's efficiency nearly in half. To eliminate the unwanted mass flux, Backhaus and Swift devised a special jet pump for the feedback loop. The strategically shaped pump creates a much larger pressure drop for one direction of motion in the oscillating gas than for the other—"like a rather leaky diode," describes Swift—and this asymmetry can be adjusted to cancel the streaming mass flow.

The output of a traveling-wave thermoacoustic engine can be tapped with an electroacoustic or other transducer. Alternatively, the engine can be used to drive an acoustic load such as a thermoacoustic refrigerator. This combination may prove quite appealing, providing refrigeration with no sliding seals, moving parts, or chlorofluorocarbons or other potentially dangerous refrigerants. "Thermoacoustic engines and refrigerators have been attractive in niche applications," comments Steven Garrett, an acoustician at Pennsylvania State University. "Now they may become even more efficient than internal combustion engines and vapor compression refrigerators."

RICHARD FITZGERALD

References

1. S. Backhaus, G. W. Swift, *Nature* **399**, 335 (1999).
2. P. H. Ceperley, *J. Acoust. Soc. Am.* **66**, 1508 (1979).
3. T. Yazaki, A. Iwata, T. Maekawa, A. Tomimaga, *Phys. Rev. Lett.* **81**, 3128 (1998).

Adiabatic Quantum Electron Pump Produces DC Current

If you know what you're doing and have suitable equipment, you can make a quantum dot, a tiny conducting region that contains anywhere from a handful to several thousand electrons. A number of experimenters have studied the electron transport properties of such dots. Now a group at Stanford University and the University of Cali-

You can control the flow of electrons in a quantum dot by cyclic changes in the wavefunction.

fornia, Santa Barbara, has made an open quantum dot, changed the shape of the system cyclically, and found that a finite current flows. This surprising effect requires quantum phase coher-