

rent gain is still very low, the intrinsic presence of rectifying Schottky diodes is an advantage for MRAM applications.

Another ferromagnet-semiconductor combination resembles a field effect transistor (FET). A FET comprises a source and a drain contact to send current through a semiconductor channel and a gate terminal to modulate the carrier density in the channel. Assume that you make a semiconductor FET with a magnetic metal for a source and a drain contact (6). The current in the channel will be magnetically polarized by the magnetic source contact. The magnetic drain picks up the current but shows a higher affinity for aligned spins. Hence, the current through the device will depend on the magnetic contact

alignment and the result is a spin switch. Furthermore, upon application of a gate voltage, the electric field perpendicular to the channel translates to an effective magnetic field for the charged particles passing through it. The effect of this magnetic field is a precession of the magnetic polarization of the channel current and for a given magnetic state of the drain, the electric current can be modulated.

The majority-minority spin situation in a metal may lead to all-metal components showing spin-switch operation (9). Metals are more suited for extreme nanoscale device sizes because of the higher number of available conduction electrons compared with semiconductors. At the nanoscale, magnetic single electron effects in Coulomb

blockade may become apparent. But first, we need to become better acquainted with the way electron spins behave in microstructured devices, of whatever material combination they are fabricated.

References

1. D. J. Monsma, R. Vlutters, J. C. Lodder, *Science* **281**, 407 (1998).
2. M. N. Baibich *et al.*, *Phys. Rev. Lett.* **61**, 2472 (1988); P. Grunberg, R. Schreiber, M. Vohl, European Patent document 0442 407(A1).
3. E. Grochowski and D. A. Thomson, *IEEE Trans. Magn.* **30**, 6 (1994); J. C. S. Kools, *ibid.* **32**, 3165 (1996).
4. *Semiconductor International* (September 1997), p. 84.
5. J. S. Moodera, I. R. Kinder, T. M. Wong, R. Messervey, *Phys. Rev. Lett.* **74**, 3273 (1995) and references therein.
6. S. Datta and B. Das, *Appl. Phys. Lett.* **56**, 665 (1989).
7. J. Heremans, *J. Phys. D Appl. Phys.* **26**, 1149 (1993).
8. G. A. Prinz, *Science* **250**, 1092 (1990); Spin-Dependent Nano-Electronics (SPIDER) ESPRIT LTR#23.307.
9. M. Johnson, *Appl. Phys. Lett.* **63**, 1435 (1993).

PERSPECTIVES: ASTROPHYSICS

Giant Planets to Brown Dwarfs: What Is in Between?

Filipe D. Santos

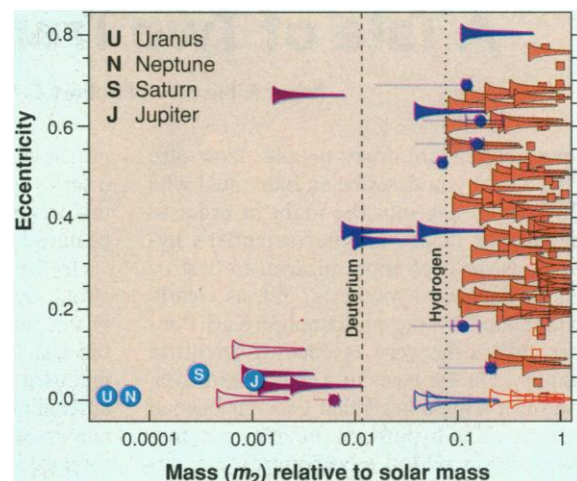
The discovery of planets orbiting nearby sunlike stars has shown that planetary systems can be surprisingly diverse and has raised many new questions about how they formed. The initial discovery in 1995 of the planet around the star 51 Pegasi by Mayor and Queloz (1) at the University of Geneva was a surprise because it is a planet with mass about that of Jupiter's and an orbital period of only 4.2 days. This implies that it is 20 times closer to its star than Earth is to the sun. The method of detection is based on the observation of periodic variations in Doppler shifts of the light from the star caused by wobbling, because of the gravitational pull of the planet. Although this method gives an accurate measurement of the object's orbital period and eccentricity, the mass deduced represents only the minimum possible mass. The Doppler effect reveals only the changes in the star's velocity along our line of sight. This means that we can only measure $m_2[\sin(i)]$, where m_2 is the mass of the planet and i is the angle of orbital inclination. The planet around 51 Pegasi has at least 0.44 Jupiter masses (M_J). Seven additional planets around solar-type stars have since been discovered with values of $m_2[\sin(i)]$ ranging from 0.44 to 6.84 M_J (2).

Enhanced online at
www.sciencemag.org

Where do we set the dividing line that distinguishes these massive planets from brown dwarfs? What are the mechanisms leading to the formation of massive planets and brown dwarfs? Brown dwarfs are frustrated stars that have insufficient mass to trigger nuclear reactions in their core. They are expected to have masses smaller than the hydrogen-burning limit of about 0.075 solar masses ($M_\odot$) but probably larger than the deuterium-burning limit of 0.013 $M_\odot$, or about 13 M_J . Companion brown dwarfs to solar-type stars have also been found by the Doppler shift method. Because of their large masses, one may attempt to detect them using astrometric measurements. This method of detection infers the presence of the companion by measuring the position of the star as it orbits the center of mass of the entire system. The advantage is that it gives the orbital inclination and therefore the real mass of the companion.

At a recent meeting on extrasolar planets, Mayor *et al.* reported that they had determined the orbital inclinations

of about half of the previously known brown dwarf candidates using precise astrometric data from the HIPPARCOS satellite (3). The figure illustrates the behavior of the orbital eccentricities versus the mass of companions of solar-type stars. In all cases of brown dwarfs that were examined by the astrometric method, the determination of the orbital inclination i resulted in a mass m_2 in the range of stars at the bottom of the main sequence, above or very close to the hydrogen-burning limit. The strong decrease in the num-



From planet to star. Objects less than 0.075 $M_\odot$ cannot burn hydrogen, and those less than 0.013 $M_\odot$ cannot burn deuterium. Companions with $m_2[\sin(i)]$ larger than 0.075 $M_\odot$ have orange symbols, candidate brown dwarfs are represented in blue, and probable giant planets are represented in red. The jovian planets are shown for comparison. The probability of having a mass m_2 larger than the minimum measured value $m_2[\sin(i)]$ is proportional to the symbol width. The objects identified by data from the HIPPARCOS satellite are represented by horizontal dotted lines from $m_2[\sin(i)]$ to the real mass m_2 [data from (3)]. The open symbols are for systems with orbits short enough to have probably been affected by tidal circularization. Circles (giant planets and objects initially thought to be brown dwarfs) or squares (stellar masses) indicate objects with known mass.

The author is at the Centro de Física Nuclear da Universidade de Lisboa, 1699 Lisboa Codex, Portugal. E-mail: fdsantos@milkyway.cii.fc.ul.pt

ber of brown dwarfs suggests that the distribution of mass of brown dwarfs does not extend to masses as small as giant planets. The new measurements indicate that brown dwarfs orbiting solar-type stars are very rare. The explanation for this rarity, although unknown at present, is probably related to the different formation mechanisms for massive planets and brown dwarfs.

Another remarkable aspect of the data is the discontinuity of orbital eccentricities of companions less massive than about $0.005 M_{\odot}$ as compared with companions in the stellar domain of masses. This behavior is in good agreement with the standard model of planetary formation. Planets are thought to originate in a protoplanetary disk of gas and dust from the collisional accumulation of successively larger planetesimals, which move in nearly circular orbits. On the other hand, initially eccentric orbits are natural in double-star systems because they result from the collapse and fragmentation processes in a mass of gas and dust.

The discovery of giant planets orbiting solar-type stars with small orbital radii raises the question of how they formed. One mechanism that has been proposed is core accretion leading to the formation of

rocky cores of about 10 Earth masses, which are then massive enough to accrete gas from the protoplanetary disk. This process requires about 10 to 20 million years to form Jupiter-mass planets. The other mechanism is gravitational instability and proceeds much more quickly, in about 100,000 years. In this process, an unstable disk breaks up into giant gaseous protoplanets where dust grains settle down. Boss proposed that the observation of optically visible young stellar objects, over a period of decades, should allow determination of which of the two mechanisms is responsible for the formation of giant planets (4). The observation of astrometric wobbles caused by Jupiter-mass protoplanets in young stellar objects with an age in the range of 0.1 to 1 million years would rule out the core accretion mechanism. However, if gravitational wobbles are found only in the older young stellar objects, the core accretion would be the favored mechanism of giant-planet formation. According to Boss, a sample on the order of 100 young stellar objects of different ages would be necessary to identify unambiguously the formation mechanism.

In both proposed mechanisms, giant planets should only form in the relatively

cool outer regions of protoplanetary disks. The discovery of Jupiter-mass planets with orbits very close to their stars causes a considerable problem because it is difficult to understand how such planets could form in place. Five Jupiter-mass planets found orbiting solar-type stars have orbital radii smaller than the distance from Mercury to the sun. The suggested explanation is that Jupiter-mass planets can form at an orbital radius of a few astronomical units and then migrate inward (5). Various migration mechanisms have been recently proposed, but it is still not possible to distinguish them observationally. Further searches with improved and diversified means of observation are strongly needed. Clearly, the discovery of planetary systems outside our solar system has opened a Pandora's box of startling phenomena and new questions.

References

1. M. Mayor and D. Queloz, *Nature* **378**, 355 (1995).
2. R. P. Butler and G. W. Marcy, *Astrophys. J.* **464**, L153 (1996).
3. M. Mayor, F. Arenou, J. L. Halbwachs, D. Queloz, S. Udry, paper presented at "Extrasolar Planets: Formation, Detection, and Modelling," Lisbon, Portugal, 27 April to 1 May 1998.
4. A. Boss, *Nature* **393**, 141 (1998).
5. N. Murray, B. Hansen, M. Holman, S. Tremaine, *Science* **279**, 69 (1998).

PERSPECTIVES: NEUROSCIENCE

A Tale of Two Transmitters

Roger A. Nicoll and Robert C. Malenka

Scientists are crazy people. How else would you describe an individual who works late into the night in order to destroy or falsify another scientist's hypothesis, or even more bizarre, to destroy his or her own hypothesis? Yet, as clearly enunciated by the philosopher Karl Popper, this is the very essence of scientific inquiry. On the basis of a few bits of data, we form a hypothesis that goes far beyond the data. The hypothesis provides a framework upon which experiments are designed to verify—or refute—the hypothesis. The longer the hypothesis can withstand these potshots, the more likely it is to be "true." More often than not, hypotheses do not withstand the onslaught of experiments and have either to be abandoned altogether or to undergo major overhauls. As cumbersome as it may seem, this is the way science advances. The history of Dale's

principle, which receives a direct hit from a series of elegant experiments reported in this issue of *Science* on page 419 (1), is a beautiful example of this process.

In the early 1930s Sir Henry Dale was struck by the strict separation of neurons in the peripheral nervous system that used the transmitter acetylcholine from those that used adrenaline (later shown to be noradrenaline). To reflect his notion that each neuron was a single biochemical unit, he proposed the terms cholinergic and adrenergic to characterize the two classes. In his 1935 Dixon Lecture (2) he expanded on this theme and developed what would later become known as Dale's principle, a modern version of which states that a neuron releases a single transmitter from all of its terminals. He suggested that the daunting task of identifying the transmitters used in the central nervous system could be eased by taking advantage of this notion—by assuming that the same transmitter is released from all of a neuron's terminals. Using the spinal primary afferents as an example, he proposed that the sub-

stance that caused cutaneous vasodilation when released by stimulated primary afferents would likely also serve as a transmitter at the afferent synapses in the spinal cord. Identification of the transmitter at one site would predict the transmitter at the other. Indeed this approach bore fruit when substance P was found to be released by these neurons (3). By this same reasoning, Eccles successfully identified the first transmitter in the central nervous system by showing that motoneuron axons, which release acetylcholine onto muscle, also release acetylcholine from their collaterals onto Renshaw cells in the spinal cord (4). Basking in the resounding success of this approach—and possibly feeling a little guilt for the heated arguments he had with Dale over the years as to whether neurons communicated electrically (Eccles) or chemically (Dale)—Eccles immediately elevated Dale's ruminations to the rarefied level of a "principle." (Although Dale always used the singular when discussing the transmitter content of a cell, he never explicitly addressed the issue of multiple transmitters in one cell; only later interpretations linked this idea to Dale's principle.)

Since its original conception, Dale's principle has undergone considerable revision. We now know that more than one transmitter can be released from a single

The authors are in the Departments of Cellular and Molecular Pharmacology, Physiology, and Psychiatry, University of California, San Francisco, CA 94143-0450, USA. E-mail: nicoll@phy.ucsf.edu