

# Silicon shines on

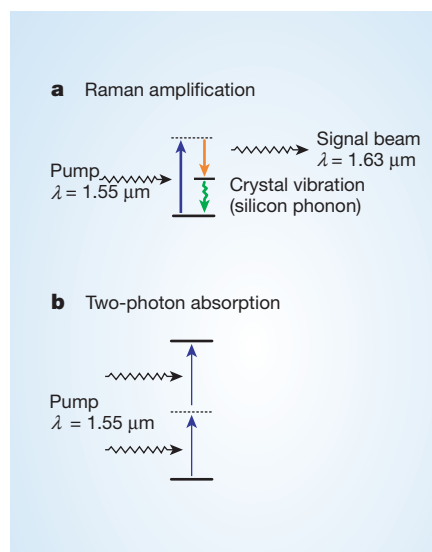
Jerome Faist

Researchers are getting better at making silicon do what it really does not like to do — emit light. A silicon laser is now demonstrated that has promising features for future practical applications.

**N**onlinear optics could be called the optical equivalent of the philosopher's stone: just as lead could be turned into gold by changing the number of protons in its atoms, so the colour of a laser beam can be changed from blue to red by crossing a nonlinear crystal. In nonlinear optics, incident light is converted to light of a different wavelength by making use of specific, nonlinear properties of a material. Recently, nonlinear optics has been found to be capable of performing another much sought-after trick — transforming silicon, the main material for electronics, into an optically active material. On page 725 of this issue, Rong *et al.*<sup>1</sup> take a further step towards a practical implementation of silicon optics by building a silicon laser that operates in a stable, continuous mode.

In the same way that steel is the base material for large-scale constructions and carbon for all known life forms, silicon is the mainstream material of electronics. It exhibits the right electronic and mechanical properties, is cheap and abundant, and can be easily processed into high-quality micrometre-scale devices. One thing that silicon could not do until recently was generate light efficiently, because of the nature of its electronic states. For this reason, all active optoelectronic applications, such as lasers for optical recording or for telecommunications, are based on group III–V materials such as GaAs and InP (ref. 2). However, as the clock frequencies of computer processors continue to increase, there is a growing need for optical data transmission that is integrated within silicon chips. Clock signals, which are needed to synchronize functions on a chip, are generated and transmitted electronically, but this scheme will run into problems with power consumption and accuracy at higher processing speeds. Optical, rather than electronic, clock distribution is expected to circumvent these problems, and the achievement of practical optical amplification in silicon would therefore be a significant advance.

The nonlinear optical effect that is used to induce light emission and amplification (laser action) in silicon is 'stimulated Raman scattering' (Fig. 1a). A laser 'pump' activates the process; a photon at the pump energy is absorbed and then re-emitted with lower energy (and so a longer wavelength) together with a 'phonon' — an elementary vibration



**Figure 1 Raman amplification.** a, In this nonlinear optical scheme, a pump photon is absorbed and re-emitted as a signal photon with a longer wavelength, along with a phonon. The process converts the pump energy into the signal beam, which is then amplified. b, Two-photon absorption, a nonlinear optical parasitic effect. It creates unwanted pairs of electrons and holes that can turn off the Raman amplification.

of the crystal<sup>3</sup>. The emitted photons make up the signal beam. By a trick of quantum wizardry, the upper energy level (dashed line in Fig. 1a) may remain virtual so that no real optical absorption is needed and the silicon crystal remains transparent. Laser action occurs because the process of light emission is stimulated — boosted — by the presence of a signal beam photon. The result is that the energy from the pump laser is transferred to the signal beam, which is then amplified.

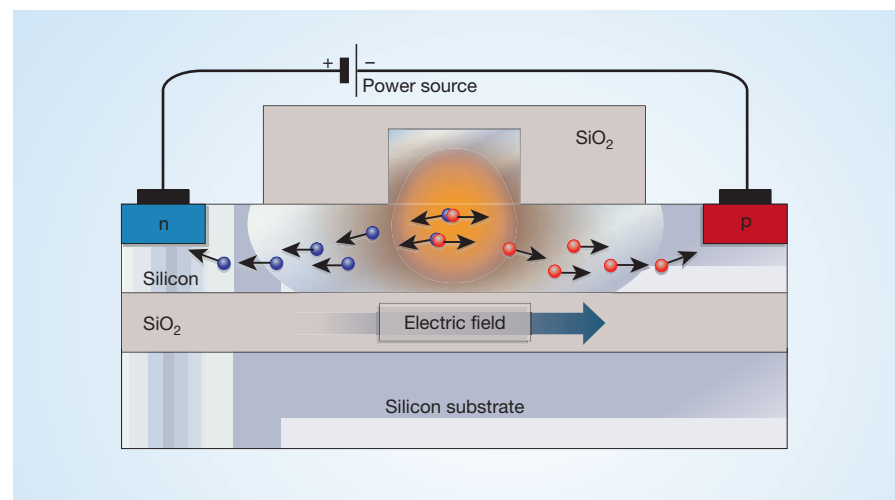
'Raman amplification' is a small effect, and to build a laser with it you need a very high pump intensity and very low absorption losses. Such conditions have already been achieved in optical devices made from silica ( $\text{SiO}_2$ ) (refs 4, 5). To achieve a sufficiently large optical intensity to produce the Raman effect in silicon, Rong *et al.*<sup>1</sup> used a recently developed silicon technology called silicon-on-insulator<sup>6</sup>. In this approach, which was originally designed to reduce the power consumption in portable electronics, thin layers of crystalline silicon with a large refractive index ( $n = 3.6$ ) are deposited on

silica layers with a low refractive index ( $n = 1.5$ ). The large step in refractive index enables a tight confinement of light, which can be exploited to achieve significant Raman amplification in silicon.

Using this approach, Rong *et al.* built a silicon waveguide structure, in the shape of a ridge, surrounded by silica to guide the light with low losses (Fig. 2, overleaf). They show that a pump laser with a power of only a fraction of a watt focused into this silicon waveguide creates an optical intensity up to  $25 \text{ MW cm}^{-2}$ , which is larger than what has been achieved within high-power semiconductor lasers. The Raman amplification at this intensity is still small (a few decibels per centimetre, compared with  $200 \text{ dB cm}^{-1}$  in standard semiconductor lasers) but is enough to produce laser action, owing to low optical losses in the silicon waveguide.

Raman amplification has already been shown in similar silicon structures, but the amplification was limited to very short pulses of a few nanoseconds at most<sup>7,8</sup>. The problem is that an unwanted nonlinear optical side effect — two-photon absorption (Fig. 1b) — creates pairs of electrons and holes that remain for a long time in the sample and absorb both the pump light and the signal light, and so quickly turn off the Raman amplification. Rong *et al.* solve this problem by embedding the silicon waveguide within a semiconductor device, a reverse-biased p-i-n junction diode (Fig. 2). This device is designed to extract electrons and holes away from the waveguide. It is formed by implanting a short region on each side of the waveguide with impurities that convert silicon into a material with electron (n-side) or hole (p-side) conduction. A positive voltage is then applied to the n-side with respect to the p-side. In this reverse-biased scheme no current flows, but a strong electric field is generated that quickly removes the electrons and holes created by the two-photon absorption effect.

With this design Rong *et al.* demonstrate a silicon laser with continuous operation, a significant advance for the development of practical silicon lasers. Of course, the use of a nonlinear optics phenomenon means that optical pumping will always be required. However, the technique converts only a small fraction of the pump power to heat in the chip, in contrast to optically pumped lasers



**Figure 2** Cross-section of the silicon laser designed by Rong *et al.*<sup>1</sup>. A ridge-shaped waveguide made of silicon is surrounded by silica (SiO<sub>2</sub>). The large difference in refractive index between silicon and silica ensures that the light intensity is tightly confined within the waveguide so that a large Raman amplification can be obtained. This structure is embedded within a semiconductor device, which enhances the laser output by draining off unwanted electrons and holes that are created by the two-photon absorption shown in Fig. 1b.

that do not rely on nonlinear optical effects. This is an important advantage given that heat dissipation is becoming the key limiting parameter in microelectronics.

A fascinating feature of this work is the use of the p-i-n junction, which combines the nonlinear-optical and semiconducting properties of silicon in the same device. Rong *et al.* show that this design enables control of the optical power emitted by the laser, which in principle should also be possible at a very high frequency and could therefore be used for information processing. Last but not least, this work demonstrates that technological advances in microelectronics, in this case the silicon-on-insulator and nanolithography techniques used to fabricate the

waveguide ridge structure, can be applied to create advances in an apparently unrelated research field such as optoelectronics.

Jerome Faist is at the Physics Institute, University of Neuchâtel, CH-2000 Neuchâtel, Switzerland.  
e-mail: jerome.faist@unine.ch

1. Rong, H. *et al.* *Nature* **433**, 725–728 (2005).
2. Rosencher, E. & Vinter, B. *Optoelectronics* (Cambridge Univ. Press, 2002).
3. Shen, Y. R. *The Principles of Nonlinear Optics* (Wiley, New York, 1984).
4. Spillane, S. M., Kippenberg, T. J. & Vahala, K. J. *Nature* **415**, 621–623 (2002).
5. Kippenberg, T. J., Spillane, S. M., Armani, D. K. & Vahala, K. J. *Optics Lett.* **29**, 1224–1227 (2004).
6. Luryi, S., Xu, J. & Zaslavsky, A. (eds) *Future Trends in Microelectronics* (Wiley-IEEE, New York, 2004).
7. Boyraz, O. & Jalali, B. *Optics Express* **12**, 5269–5273 (2004).
8. Rong, H. *et al.* *Nature* **433**, 292–294 (2005).

## Cell biology

# Divide and conquer

Michael Hengartner

The discovery that cell death in nematode worms induces fragmentation of mitochondria reveals a new parallel to the death process in mammals, and may shed light on why mitochondria divide in death.

When mammalian cells die by the process of apoptosis, their mitochondria fragment into smaller pieces. Why these power-generating compartments should divide as the cell around them dies, and whether this fragmentation is important for the death process or simply an epiphenomenon, has so far largely remained unclear. But an answer is suggested by the paper from Conradt and colleagues on page 754 of this issue<sup>1</sup>. The authors show that mitochondria also fragment during apoptosis in the small nematode worm *Caenorhabditis*

*elegans*. Moreover, experimental induction or prevention of mitochondrial fragmentation could respectively enhance or partially prevent apoptosis. These observations hint that mitochondrial fragmentation has an evolutionarily conserved, causative role in promoting apoptotic cell death.

The term apoptosis refers to a specific type of programmed cell death that occurs in all multicellular animals, from the lowly worm to the highly complex human. Apoptosis is characterized by specific morphological changes in the dying cell,

and is mediated by several protein families<sup>2</sup>.

In mammals, most apoptotic cell deaths are mediated by a specific signalling pathway known as the mitochondrial pathway. As its name implies, this pathway requires the active participation of mitochondria — the organelles better known for their role in cellular respiration and the generation of the high-energy molecule ATP. In cells condemned to die, mitochondria release several dozen proteins into the cytosol, and they can then wreak havoc in the rest of the cell.

The best known of these mitochondrial expatriates — cytochrome *c* — interacts in the cytosol with the Apaf-1 protein, ultimately activating a group of proteases (protein-digesting enzymes) known as caspases<sup>3</sup>. These enzymes then cleave a selected set of target proteins, resulting in the controlled ‘implosion’ of the cell. How cytochrome *c* *et al.* manage to cross the outer lipid bilayer of the mitochondria to reach the cytosol is still hotly debated. What is clear, however, is that this release is regulated by proteins of the Bcl-2 family, many of which can bind directly to the outer mitochondrial membrane.

Recently, several groups have reported a second peculiar behaviour of mitochondria during apoptosis: not only do they release proteins, but they also fragment into smaller pieces<sup>4</sup>. That mitochondria can fragment is nothing new in itself — like bacteria, mitochondria divide by a process of fission, in which one long organelle is pinched in the middle to produce two shorter daughters. Unlike bacteria, mitochondria can also undergo the reverse process, and fuse together to form long filaments. Fission and fusion are tightly controlled, and are important for the proper distribution of mitochondria during cell division.

But why mitochondria should fragment during apoptosis is not clear. One possibility is that the release of mitochondrial proteins stimulates mitochondrial division. Indeed, conditions that compromise mitochondrial function have been reported to result in short, round mitochondria. An attractive alternative would be that fission is necessary (directly or indirectly) for the release of cytochrome *c*. Consistent with this idea, interfering with the fission process has been reported to delay cytochrome *c* release during apoptosis<sup>5</sup>. Whether this is a general phenomenon remains to be seen. Furthermore, given that mitochondrial fission occurs continuously in living cells, there must be more to the story than fission simply promoting death.

Enter *C. elegans*. Genetic studies in this species showed that most components of the apoptotic pathway have been conserved throughout evolution<sup>6</sup>. For example, *C. elegans* has a Bcl-2 counterpart (CED-9), an Apaf-1-like molecule (CED-4) and a caspase (CED-3). Surprisingly, however, mitochondrial proteins have so far played at best a minor role in the apoptosis saga in *C. elegans*.