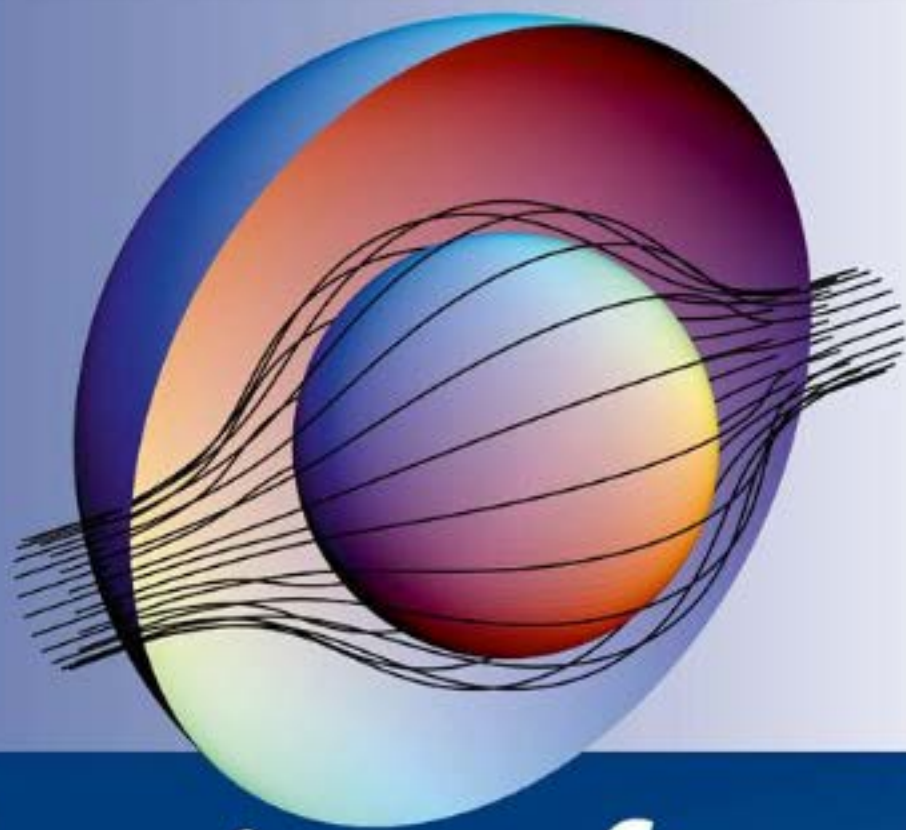


Martin Beech



The Physics of Invisibility

A Story of Light and Deception



 Springer

The Physics of Invisibility

Martin Beech

The Physics of Invisibility

A Story of Light and Deception

 Springer

Martin Beech
University of Regina
Regina, SK, Canada
martin.beech@uregina.ca

ISBN 978-1-4614-0615-0 e-ISBN 978-1-4614-0616-7
DOI 10.1007/978-1-4614-0616-7
Springer New York Dordrecht Heidelberg London

Library of Congress Control Number: 2011938685

© Springer Science+Business Media, LLC 2012

All rights reserved. This work may not be translated or copied in whole or in part without the written permission of the publisher (Springer Science+Business Media, LLC, 233 Spring Street, New York, NY 10013, USA), except for brief excerpts in connection with reviews or scholarly analysis. Use in connection with any form of information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed is forbidden.

The use in this publication of trade names, trademarks, service marks, and similar terms, even if they are not identified as such, is not to be taken as an expression of opinion as to whether or not they are subject to proprietary rights.

Printed on acid-free paper

Springer is part of Springer Science+Business Media (www.springer.com)

*This book is for Gael with many thanks
for showing me the stars from New Zealand.*

About the Author

Martin Beech is Professor of Astronomy at Campion College, The University of Regina in Saskatchewan, Canada. He has written numerous research articles on topics ranging from meteor physics, Martian meteorites, stellar structure and evolution, cosmology, the history of science and mathematical number theory. He lives in Regina with his more than visible wife, a brother-in-law, five dogs (*a.k.a.* The Five Pugs of the Apocalypse) and three cats – there are times when he certainly wishes he could be invisible.

Introduction

The second worst joke in the world begins with the question, “When is a door not a door?” The answer to this conundrum, of course, is, “when it is ajar.” The logician might at first wince at this dreadful play on words, but then in learned tones continue that a door is not a door, of course, if one cannot see it in the first place, or if it is camouflaged to look like something else.

Well, the joke, as we first warned, is the second worst joke in the world, and the logician’s rejoinder is no more amusing, but at least there can be little argument with the fact that if you can’t see something then it is impossible to name it, and if it looks like something else then the naming will be wrong. For, indeed, we live in a world of names and light. Through an incredible act of evolution our eyes bring the light from the surrounding world to the rods and cones attached to our retinas, and there a photochemical reaction transforms the laid-out vista, a real and inverted vista at that, into a series of electrochemical pulses that our brain can, somehow, transform into a mind’s-eye view. Reality turned into light, into chemistry, into synaptic pulses, into virtual inner seeing – what an incredible transformation our view of the world truly is.

It is by this remarkable chain of transformations that we see and thereby recognize and name the objects around us. Indeed, the world is given permanence by naming: a door is a door and nothing else – unless, of course, it is ajar, or if it cannot be seen. We make sense of the world and its stability by seeing and interpreting the images that are presented to us, and we take dominion over the world by giving names to those objects that don’t change their form from one moment to the next. By the act of seeing, and through the action of naming, the world makes sense to us – well, some of the time. And yet, the Belgian surrealist painter René Magritte also reminds us of “The Treachery of Images” (Fig. 1) and how the mind’s-eye can be so easily fooled. Not least with respect to context can we be hoodwinked, but our perception of reality can also be misdirected and manipulated – what we think we see is not necessarily what occurred or was actually there. Indeed nature, long ago, learned that survival is all about blending into the background and looking like something else. Camouflage, after all, is about the art of melding into invisibility. So let’s ask the question again: “When is a door not a door?” Answer: When it looks



Fig. 1 "The Treachery of Images," by René Magritte, oil on canvas painted 1928/29. This picture does not show a pipe. What we see in this image is, in fact, the image of an image of a painting of an object called a pipe



Fig. 2 *Cnemaspis Neangthyi* – a new species of gecko discovered in the rocky foothills of the Cadamon Mountains of Cambodia in 2007. The gecko's surface texture and coloring is a superb camouflage match to the background rock. Image courtesy of Lee Grismer, La Sierra University, California

like a brick wall. Likewise, "When does a gecko not look like food?" Answer: When it looks like a piece of rock (Fig. 2).

Without light all is lost. If we cannot see an object we cannot name it, and it is invisible to our senses – it becomes a mystery and at best a shadow-form devoid of meaning. For all this, however, the world is full of impossible images; not least among these are the shimmering mirages on a deep blue sky, turning distant sailing



Fig. 3 Two of the very first experimental cloaking devices. These disks can mask from view (at microwave wavelengths) any object placed within their central core (Image courtesy of D. Smith, Duke University)

ships into airborne galleons, or the diaphanous glow of Pepper’s Ghost, the phantasm that haunts many a thespian stage. Reality, for so it would seem, is mercurial and ever changing, and what we think we see may be nothing more than an imagined chimera – as French author and diarist Anaïs Nin so aptly wrote, “We don’t see things as they are, we see things as we are.”

There are many, many ways in which an object can be made invisible, and some of them, as we shall see, appear to fall under the banner of magical. This is the magic of science, of course, rather than the magic of wishful thinking or the miraculous act in which something happens beyond the laws of physics and normality. Science is, indeed, a hard task master, but fortunately it has not stifled the free ranging nature of literary imagination – for certainly, both ancient and modern writings are rich with the ideas of magic, invisibility and the miraculous. The desire to walk unobserved through the world and to command invisibility at will is a power that many humans may have wished for, but as Chap. 1 reveals it is a power that no one person should have unrestricted access to. But this being said, we now live in a world where the technology to make solid objects invisible to outside detection does exist. It is a reality, and we assuredly do live in an age of miracles and wonders (Fig. 3). These fledgling devices will change the world and the way in which we see it – or don’t see it, as the case may be.

The world is revealed to us by seeing, and there is both fear and desire in the thought of invisibility. It would, indeed, be the strangest sight we ever saw if something disappeared before our eyes. And yet, our waking world is full of familiar invisible entities. The flow of electricity along a bare metal wire, the gravity that keeps us grounded, the air we breathe, the bacteria and viruses that make us ill, the X-rays that reveal our broken bones, the time that measures our days, and the radio signals that keep us connected – all are invisible to our visual senses, and yet they surround us and we don't give them a second thought.

Not only this, we also explain the very world in terms of entities we cannot see. The atoms that make up this page and the chair you are sitting on – even you – are so minutely small that to our everyday senses they are not just invisible, they verge on the non-existent. Indeed, the famous Austrian physicist and logical positivist philosopher Ernst Mach (as recently as a century ago) denied the very existence of atoms because no one could observe them. And yet, the modern theory of the atom is one of the most successful and most cherished of all scientific accomplishments.

Not only do we humans give names and permanency to those objects that can be seen, we also provide names and existence to entities that we cannot sense at all. Invisible doesn't necessarily mean out of sight and out of mind. "Vision," in the intellectual sense, Jonathan Swift reminds us, "is the art of seeing what is invisible to others."

To understand what it is that we see, we must first understand light itself, and in this era of nascent cloaking devices this also means that we must know how to control light and make it move according to our desires. But what is light? This is one of the questions that will be answered in Chap. 2. Before this, however, in Chap. 1 our attention will be directed towards a brief review of what invisibility has historically meant to society, scientists and philosophers, and we shall look at a few of the simpler invisibility-making ruses. In Chap. 3 we begin to explore some of the ways in which light can be manipulated and how some more complex forms of invisibility are achieved. Chapter 4 delves deeper into the question of light as a manifestation of electromagnetic radiation, while Chap. 5 turns to issues relating to the direct manipulation of light via metamaterial substrates (Fig. 3).

Much of the physics of invisibility seems magical and often downright impossible, but in the world of scientific exploration there are many different levels of magic and impossibility. Perhaps, however, as physicist Michio Kaku has explained in his highly entertaining book *The Physics of the Impossible* (Anchor Books, 2009), there are three basic levels of impossibility. Type I impossibilities encompass those ideas that are presently impossible to realize, but the way to make them possible is at least known in principle – in other words they do not contradict any known physical principles. Type II impossibilities are those that are currently impossible, and it is not clear how they might be made possible in the future either technically or theoretically. And, finally Type III impossibilities constitute those processes that are impossible because they violate one or more of the known laws of physics – so-called perpetual motion machines fall under this category, for example, because they violate the known (and fully agreed upon) laws of thermodynamics.

Making an object invisible to all wavelengths of light is presently a Type I impossibility – since, as we shall see in later chapters, it is already possible to make an object invisible at selected wavelengths of electromagnetic radiation. The final construction of a fully flexible, all wavelengths of light invisibility cloak is perhaps a few decades away. At issue are not fundamental questions of theoretical physics but rather questions relating to technology development and fabrication methods. For all this, however, there are many ways of making something disappear, and a few of the more recent ideas hover tantalizingly in the borderline realm separating Type II impossibilities from those of Type III. The fun, of course, will be in finding out exactly which side of the borderline they do reside upon. The physics of invisibility is a new, dynamic and exciting field of research, and we shall surely marvel at what we may see, or more appropriately not see, as the near-term future unfolds.

Contents

1	Of All Things Visible and Invisible.....	1
	Ancient Mythology	2
	Unseen Beasts	3
	Magical Powers.....	4
	Invisible Ink	6
	Smokescreen	7
	Dimension Shift	9
	The LHC and Extra Dimensions.....	11
	Irreversibly Invisible: Black Holes	12
	Time Shift	14
	In the Eye of the Beholder	14
	The Cottingley Fairies.....	16
	The Multiverse	17
	The Trojan Horse	19
	The Vanishing Pyrex Block	19
	Wood and Strutt Explain	20
	Blending In.....	21
	Seeing the Impossible	23
	The Ultimate Optical Illusion	27
	Tripping the Light Fantastic.....	28
2	A Very Brief History of Light	29
	Beginning at the Beginning.....	29
	The Act of Seeing	30
	The Laws of Reflection and Refraction	32
	The Writings of Alhazen.....	33
	The Dreams of Dr. Mirabilis.....	34
	Too Quick, Even for Galileo	35
	Hooke’s <i>Micrographia</i>	36
	Moving Heaven and Earth	37

Dealing with Descartes	38
Snell's Law.....	38
Fermat Finds the Quickest Way	40
Newton's <i>Opticks</i>	42
The Inconstant Moons of Jupiter	43
The Aberration of James Bradley	45
Waves and Wave Interactions.....	46
The Triumph of Waves.....	47
Faster or Slower? A Crucial Test	49
Polarization	52
Sight and Sound	53
3 Mirror Tricks, the Art of Cloaking and Seeing the Impossible	57
Smoke and Mirrors: Without the Smoke	57
The Infinity Box	58
Pepper's Ghost	59
The Tachi Invisibility Cloak.....	60
A Moment's Reflection	63
Cats' Eyes	64
Seeing Through a Door.....	67
Alpha Blending to Half Transparency	69
Crystal Vision and the Invisible Touch	70
Augmented Reality	72
Fata Morgana	73
Earth's Atmosphere and the Sunset Deception.....	75
The Green Flash.....	76
How to Catch an Atom Bomb.....	77
Mirages: Inferior and Superior.....	78
A Cosmic Mirage	80
Proving Einstein Right.....	82
The Impossible Sunrise.....	85
Optical Fibers and Deep-Sea Sponges.....	87
From High Fashion to Cuttlefish and Squids	89
Eaton's Singular Cat's Eye Lens.....	91
Maxwell's Super Fisheye Lens	93
The Luneburg Cloaking Array	94
A Crystal Invisibility Cloak	97
The Cone of Silence.....	99
Sonic Crystals	100
A Wave-Less Sea	102
And Earth Stood Still	104
4 Maxwell's Marvelous Waves	107
Magnetic Attraction	107
Laying on the Leyden Jar.....	108
Michael Faraday.....	111

Faraday's Lightning Cloak.....	115
Gauss Lays Down the Law	115
A Beginning with String	116
A Prelude to Electromagnetic Waves.....	118
Making Waves.....	119
All Waves Together.....	123
Lights Out for the Aether.....	123
Putting the Particles Back	126
Absolute Light Speed.....	127
Warp Speed and Absolute Invisibility Cloaking.....	127
5 The Metamaterial Revolution	131
The Black Arts	131
The Strange Paper by Dr. Veselago.....	133
Marconi All at Sea	139
Negative Refractive Index Engineered.....	141
Split Ring Resonators and Metamaterials.....	142
The Harry Potter Project	146
Superlenses	148
The First Vanishing	153
Nanomaterials and a Small Slice of Black.....	156
Total Absorption	158
Carpet Cloaking	161
Nothing Is Perfect	171
Transformation Optics	174
Harry Potter's Sleeves and Eyesight for the Blind.....	179
Where Next for Metamaterials?.....	180
Optical Dynamics and Experimental Cosmology.....	181
Out of Space and Time.....	185
Afterword.....	189
Appendix 1 Black Holes in the Laboratory	191
Appendix 2 Atmospheric Refraction, Extinction and Scattering.....	195
Appendix 3 How Big Is a Photon?.....	199
Index.....	201

Chapter 1

Of All Things Visible and Invisible

*O world invisible, we view thee,
O world intangible, we touch thee,
O world unknowable, we know thee,
Inapprehensible, we clutch thee.*

Francis Thomson (1859–1907)

It is incredible. Out of all the epochs of human history we find ourselves alive in the era in which magic is truly manifest. This wizardry is no hocus-pocus, however, but resides in the marvelous new technologies that can make a solid object disappear from view. That such knowledge and engineering skills exist is a triumph of human ingenuity, but, for all this, we must ask of this new innovation, where will it take us? Indeed, the desire to control our visibility has deep and dark roots that stretch back to the very beginnings of human antiquity. Some reflection seems entirely appropriate, therefore.

How, for example, did the idea of controlled invisibility inspire our distant ancestors, and how have such ancient dreams permeated our modern culture? Indeed, through numerous novels, plays, films, and more recently television programs, invisibility, in its many different guises, has been exploited as a fruitful plot theme. For invisibility has been, and continues to be, a rich literary vein to examine and explore. It has it all; ranging from the philosophical, the historical, the ethical, the artistic, the mythical, and sometimes the downright bizarre. Indeed, as we shall see, the history of invisibility is a strange and wonderful one, and it has, on occasion, preoccupied the minds of some of humanity's greatest thinkers.

We live in the corporeal world; the things we see and sense have substance and form. We need these perceptions of reality to function, and yet for all this the desire to become invisible, at one's own command, is a craving as ancient as humanity is old. We can all probably think of times when we have wished to be unseen, to be a fly on the wall – an invisible interloper with the ability to eavesdrop and secretly see what others are up to. And yet, the power of invisibility is a double-edged sword. It is both a blessing and a curse. To the lonely of heart, to see and be unseen is an

unbearable grief, but to every spy and despot it is a wanton desire. Indeed, throughout recorded history the pitfalls and ethical challenges that would attach to the gift of becoming invisible have been debated and written about. Let us see, therefore, some of these ancient and venerable musings.

Ancient Mythology

Rising from the abyssal depths of legend we first encounter controllable invisibility through the Helm of Haydes. Wrought by subtle craft and skillful hand this helm was the secret weapon wielded by the Lord of the Underworld. Indeed, it proved of critical advantage in the Titanomachy – the divine war between the Olympian and Titan gods of old. The Helm had the power to make its wearer invisible, and Haydes used it to slip into the Titan camp and thereby destroy their weapons, ushering in the victory of the Olympian gods. Invisibility, it would seem, is not just the desire of mortals but also the weapon of ancient deities.

A magical cap, *tarnhelm* (literally meaning magic helmet), also appears in Norse mythology, and Wagner's orchestration of the *Ring of the Nibelungen* includes reference to its ability to make the wearer invisible. More than this, however, *tarnhelm* can also work to change the form of the wearer as well as transport them to distant places instantly – a powerful device indeed. J. R. R. Tolkien in his famous tale *Lord of the Rings* also makes use of such a magical helmet, via the English equivalent name of *dermhelm*. In this case the shield-maiden Éowyn adopts the name of the helm in her disguise as a male warrior. In Tolkien's application the wearer is not made physically invisible, but their outward persona is transformed.

More than simply existing as a transformative agent, however, there is a cost attached to the employment of an invisible cloak. Indeed, unfettered access to such a device will ultimately warp the moral compass of even the most innocent of users. The Greek philosopher Plato understood this well, and in Chap. 2 of his masterpiece *Republic*, written circa 370 B.C., he recounts the mythical story of the Ring of Gyges – a magic ring that gave the wearer the ability to become invisible. The story concerns a lowly shepherd who, after an intense rainstorm, finds that an underground chasm has opened up close to the location of his flock. Entering the chasm he discovers inside a bronzed horse statue. Cut into the side of the hollowed-out horse are windows, and within the statue's interior is the naked body of a giant dead man. The strange corpse has but one adornment, and that is a golden ring, which the shepherd, perhaps not unsurprisingly, takes.

Purely by accident the shepherd discovers the power of the ring, and he immediately thereafter arranges with his villagers to become a delegate to the king's palace. The corruption of the ring appears to act immediately upon the shepherd, who once inside the palace sets about seducing the king's wife. The shepherd then kills the king and claims the throne for himself. It is a veritable rags to riches story for the shepherd, but one that costs him his humanity. Indeed, this is undoubtedly the idea that Plato intends the story to portray. In Robin Waterfield's recent translation

of the *Republic*, we read “There is no one... who is iron-willed enough to maintain his morality and find strength of purpose to keep his hands off what doesn’t belong to him, when he is able to take whatever he wants... without fear of being discovered.” Plato’s story of the Ring of Gyges is short, taking about a page of text to tell, but it is full of deep-searching questions. How, indeed, would anyone, whether from the ancient past or the present day, behave if they had the power to make themselves invisible?

Where are the limits beyond which no invisible mortal should go? For Plato, there appeared to be no limits. Eventually any device that could render its owner invisible upon command would fully corrupt that owner, sending him or her into a descending spiral of moral defeat and debauched behavior. The temptations would be just too great, and no one would be able to stave off the inner conflict for very long. At heart this is also the message in J. R. R. Tolkien’s *Lord of the Rings*. The Ring, with its associated ability to turn its wearer invisible, is ultimately overpowering, all-consuming and morally corrupting. Indeed, no long-term good can come from it at all. It is a power that no human (or elf or dwarf) should ever have unfettered access to.

Unseen Beasts

While invisibility will bestow great power on the user, it will automatically instill great fear into those who lack such ability. Writers and film makers have long realized this latter fear, and at the core of many a horror story is an unseen, shape-shifting, invisible monster. Indeed, we fear the unseen and even the dark – because, of course, when it is dark our visible senses fail us, and we are blind to our surroundings.

The thought of being stalked by unseen monsters intent upon our destruction is universally disturbing. The invisibility in this case gnaws at our innermost, primordial fears. This was the very premise, for example, in the climactic scenes of the movie *The Silence of the Lambs* (Orion Pictures, 1991). Here the hero FBI cadet investigator is stalked in the darkness, deprived of sight, and left facing an invisible enemy who uses night vision goggles to stalk his victim. Ultimately it is sound that gives the location of the invisible villain away. Here, in fact, is one of the enduring problems of practical invisibility: how to see your surroundings when you can’t actually be seen and how to remain fully undetectable. A similar story line is developed in the classic science fiction film *Forbidden Planet* (MGM, 1956), where the invisible manifestation of a long-extinct civilization’s consciousness terrorizes a team of investigators on planet Altair IV. The unseen monster, the Id, is eventually revealed when it tries to move through a defensive force field.

Ex-U.S. military chemist turned fantasy writer Sterner St. Paul Meek (1894–1972) highlighted this same invisibility problem in his short story *The Cave of Horror*. Published in the January 1930 issue of *Astounding Stories of Super-Science*, this horror tale concerns the discovery of a giant but invisible frog (ah, the days of innocent monsters) in the Mammoth Caves of Kentucky. While invisible to human eyes

the monstrous beast can be revealed through ultraviolet light. Having thus found the weak link in the oversized amphibian's invisibility armor, the hero of the story, Dr. Bird, is able to photograph the demon with a camera using quartz, rather than ordinary glass, lenses. The latter would absorb UV light rather than transmit it and bring it to a focus. Making an object invisible at all wavelengths of light remains, in fact, one of the biggest technical challenges of present-day research, and becoming fully invisible is, as we shall see, far from easy.

British author Herbert G. Wells explored the problems of isolation and madness in his now famous 1887 serialized novel *The Invisible Man* (Fig. 1.1). The book, although a science fiction classic, is a rather weak literary affair. The subtitle of the text, *A Grotesque Romance*, partly gives the storyline away, but tragic farce, as suggested by reviewer and novelist Arnold Bennett, is closer to the truth. Nonetheless, Wells explains in his narrative that the invisible man, simply introduced as Griffin, was once a prize-winning chemistry student who became fascinated by light – a fascination that ultimately resulted in his wretched death. Interestingly, Wells envisioned and described Griffin as being a 6-ft-tall albino with red eyes set into a pink and white face. By becoming invisible, therefore, Griffin is transformed from a person who would automatically stand out in any crowd to one that was entirely unnoticeable.

Griffin explains in the text that “I found a general principle of pigments and refraction – a formula, a geometrical expression involving four dimensions...without changing any other property of matter – except in some instances, colors – to lower the refractive index of a substance, solid or liquid, to that of air....if it neither reflects nor refracts or absorbs light, it cannot of itself be visible.” The terms refractive index, reflection, and refraction will be explained in Chap. 2, but the point is that Wells is absolutely correct in a physics sense – if, of course, one could reduce the optical properties of a substance to that of air (an experiment that illustrates Wells's idea will, in fact, be described below).

Although the physical principle invoked by Wells is sound enough, he thought that it was doomed to failure in reality. Writing to Arnold Bennett shortly after the book first came out, Wells commented that he believed the practical problems of invisibility to be “insurmountable,” but as with many things that Wells predicted (atomic bombs, space travel, and radio communication), science has begun to catch up with his writing – and in this case proved him wrong. Indeed, one present-day researcher in the field of active camouflage techniques, Masahiko Inami at the University of Tokyo, has commented that “making objects virtually transparent, like an H. G. Wells *Invisible Man*, is a [long-held] human dream.”

Magical Powers

Although many authors have reflected upon the ethical and moral issues, invisibility has also been used, on numerous occasions, to simply move a storyline along. The invisibility in this case may provide a given hero with the morally benign ability to

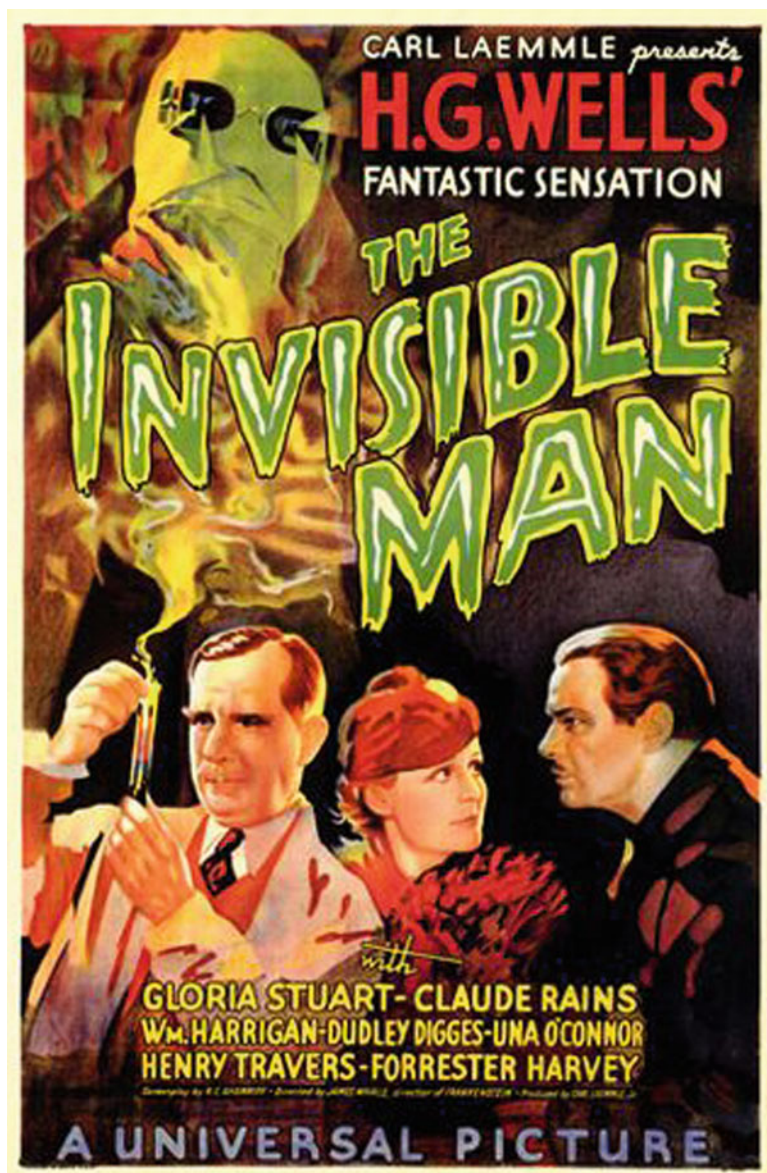


Fig. 1.1 Advertising poster for the 1933 Universal Pictures film version of *The Invisible Man* directed by James Whale

learn some important plot detail or otherwise hidden secret. Although magic and its implied devil worship was very much a taboo, if not a downright dangerous subject at the time, the great playwrights of Elizabethan England knew that such themes were both crowd-pleasers and a useful means of plot development. Christopher Marlowe

(1564–1593), for example, used the ruse of character invisibility in his play *The Tragical History of Doctor Faustus* (first performed on stage circa 1592 and published posthumously in 1604). In this particular case, Mephistophilis, the devil's messenger and gift-granting guide to the overambitious Dr. Faustus, induces invisibility so that they can make a secret visit to the Vatican, where Faustus plays mischievous games on the monks and hurls insults at the pope.

Although Marlowe invokes invisibility through the art of black magic, his contemporary, William Shakespeare (1564–1616), used the same idea in a slightly more benign, folklorish manner. In *The Tempest* (written circa 1610) Shakespeare enables Ariel, the spirit-servant rescued by the magician Prospero, to move invisibly “like a nymph o’ the sea.” Ariel, essentially, becomes Prospero’s disembodied eyes and ears, moving, as he does, unseen between the various travelers and sailors that have been deliberately shipwrecked upon the magician’s island. Although invisibility via magical manipulation is invoked in both Marlowe’s *Faustus* and Shakespeare’s *Tempest* it is ultimately used in very different ways. In *Faustus* it is applied in an entirely selfish and debauched manner, while in the *Tempest* it is employed in a basically harmless sense, essentially allowing plot information to be obtained by the characters.

In more recent times, J. K. Rowling, author of the highly successful Harry Potter series of books, developed a number of themes using invisibility. Indeed the cloak of invisibility plays a central role in the story. This cloak enables the young hero to move unseen around the corridors of Hogwarts school and to even launch an attack on his arch enemy Draco Malfoy. Although acting to defend his friends in this latter instant, a laudable action, Harry is nonetheless in danger of falling under the morally corrupting influence of the cloak by attacking as an unseen agent.

Rowling is aware of the dangers posed by allowing Harry, the morally good character of the story, to have unfettered access to an invisibility cloak, and indeed, we are later informed in the narrative that it is one of the three deathly hallows. The cloak, in fact, has the ability to hide its wearer from the otherwise withering gaze of the Grim Reaper. Here is power indeed. Interestingly, Rowling also indicates that the invisibility cloak will not hide its wearer from all forms of detection. The ghostly Dementors, for example, would not be fooled by the invisibility cloak since they sense their victims through emotional contact. Likewise, anyone wearing the cloak would still be detectable through their body heat or through the sense of motion. These very same issues are some of the key areas of modern research with respect to cloaking devices. Somebody or something may well be invisible to sight, but there is more than one way of sensing a presence.

Invisible Ink

In *Harry Potter and the Chamber of Secrets* (published in 1998) Rowling uses perhaps the most familiar, almost everyday, example of actual invisibility – that of invisible writing. In this case a message is hidden from the readers’ view until

revealed by the appropriate viewing procedure. Even when written in invisible ink, however, one can never be sure who else might eventually read the message.

The idea of writing an invisible message at first sounds like an oxymoron, but as all experts in espionage well know it is a vital skill required of every secret agent and spy. There are many varieties of invisible ink, and they work in many different ways. Some use chemical reactions to reveal the message, others use heat and yet others use fluorescence under certain lighting conditions.

The simplest forms of invisible ink exploit the properties of organic compounds to oxidize and thereby turn brown when gently heated. Substances such as soapy water, milk, and lemon juice work nicely and are a cheap means of making invisible ink. The message once written and allowed to dry is invisible under normal lighting conditions but is revealed once the paper is gently heated over, say, an electric light bulb. If you want your message to appear purple when resolved then you would write with a dilute solution of cornstarch (mix 1 cornstarch measure to 2 of water). By lightly running a sponge dipped in dilute iodine (mix 1 iodine measure to 10 of water) the colored text will appear. A slightly more sophisticated invisible ink message can be written with some specific kinds of laundry detergent. The key here is to use a detergent containing what are called optical brighteners – the back of the product package will indicate if brighteners are present. The message once written and allowed to dry will be invisible under normal lighting conditions but can be revealed under a black (that is, ultraviolet) light source. In this case the UV light causes the optical brighteners to fluoresce, thus revealing the message.

A more high tech example of invisible writing is that found on the back of most major credit cards. As a security device the tape upon which the signature of the card holder is written will glow, when illuminated by UV light, with the credit company's name.

References to the use of invisible inks of one kind or another stretch back thousands of years, and they have seen much use in the military services. Indeed, George Washington used invisible ink to report on British activities during the American Revolution, and during the trench fighting of World War I messages written with invisible ink were used for the transmission of orders. The inherent problem with such messaging, however, is that it is not always possible to tell if the message has been intercepted and read by the enemy.

Smokescreen

While literary scholars and philosophers have tended to focus upon the morally corrupting and emotional risks of invisibility, military strategists and nature have fully embraced it for its very definite survival possibilities. From a military perspective the ability to make troops, aircraft, and field weapons invisible is of great strategic value; so, too, is the ability to approach the enemy unseen. Indeed, it is arguably unethical to wield such power, since the ability to remain invisible is hardly fighting fair. Of course, war has nothing to do with fairness, and it is hardly surprising, therefore,

that virtually all of the current research on invisibility and cloaking devices is being provided for by military funding agencies.

In times of conflict, invisibility can be achieved by either merging seamlessly into one's surroundings, or by changing the surroundings to hide any important activity from external view. To this latter end, armed forces throughout history have deployed smokescreens to conceal their operations from view, and while the smoke itself is very obvious to the enemy, the tactical point is that it readily shields troop and vehicle movements. In the modern era special forms of smoke are engineered to not only mask any activity at optical wavelengths but also in the infrared, at which wavelength night-sights operate. There is even a super-dense smoke that shields activities from laser ranging and missile guidance systems. In a strange Janus-face sense, smokescreens, by their very nature, are plainly obvious and yet still function as an invisibility cloak.

Smoke is a complex substance and is composed of a colloidal gas containing particles and fluid droplets that interact with and scatter light rays. In reality a smokescreen doesn't actually block an external observer's view of what's happening behind it but rather scrambles the image – scattering the light so much as it passes through the miasma of particles that the original coherent image becomes completely garbled. The same effect happens when the weather is foggy.

The scattering of light is described by what is known as Mie theory, after the German physicist Gustav Mie, who first worked out the details in the early 1900s. The scattering is actually related to light rays interacting with the electrically charged component of the atoms making up the fog (or smoke, as the case may be), and it is also sensitive to the physical size of the constituent particles of fluid droplets. A form of frozen fog is mentioned by Antarctic explorer Apsley Cherry-Garrard in his wonderful account of *The Worst Journey in the World*. Describing the view over McMurdo Sound in March 1911, he writes, "It is difficult to see long distances across open water at this time of year because the comparatively warm water throws up into the air a fog, known as frost-smoke.... If calm the smoke rises and forms a dense curtain."

Cherry-Garrard's use of the term curtain is apt, since some research groups are presently looking at ways in which very small, indeed nanoparticle structures might be added to water pools or mists to produce a camouflage screen that can be actively controlled. The idea behind such screens is to hide equipment or troops from view by presenting the external viewer with a benign image. In this way their attention is not drawn to the scene, since it presents nothing out of the ordinary, and the activity behind it thereby remains undetected.

A parallel idea is also under development by weapons' manufacturer BAE Systems, where military tanks are being given the ability to "sweat." In this case the sweat, or water sheen, on the tank's surface will help to reduce its infrared heat signature, thereby making it more difficult to detect and target. In addition, by varying the flow rate and location of the water streams across the tank's surface, it can be made to mimic the shape of almost any object at infrared wavelengths, and the deadly tank can be transformed into a harmless looking building, if stationary, or transport truck, say, if moving.

Dimension Shift

Smokescreens and camouflage are all practical ways of hiding something from view. There are, however, many possible ways in which an object might in principle be made invisible, and a few such methods are described below. These examples, however, fall under the heading of theoretically imaginable but not necessarily doable. Indeed, they are ideas that currently straddle the fine, razor-edge separating Type II impossibilities from those of Type III.

One such – in principle – way of achieving invisibility, commonly encountered in science fiction and fantasy literature, is to literally move out of space and on occasion time. The idea here is to shift to a dimensional realm beyond the three (up–down, left–right, and backwards–forwards) that we normally perceive. Mathematicians and physicists have long used the concept of multiple-dimensional space in their studies, and while it may not be possible for us to perceive these extra dimensions, they nonetheless have some associated reality and accordingly offer a potential place for hiding secrets.

Four-dimensional space is the simplest extension of our normally perceived three-dimensional surroundings, and in the same way that we can think of a three-dimensional object casting a two-dimensional shadow, so the shadow cast by a four-dimensional object is a three-dimensional construct. In principle, then, if we can somehow utilize and manipulate a fourth spatial dimension, we might be able to weave the shadow that it casts to hide an actual three-dimensional construct. What we would see (the three-dimensional shadow shape) would not, in this manipulated situation, correspond to the actual object located in front of us. In this case invisibility is achieved by a visual transformation effect – it is a kind of Trojan horse (a topic to be discussed later).

Interestingly, the idea of building a four-dimensional, so-called hypercube house was explored over 70 years ago by the American science fiction writer Robert A. Heinlein in a short story published in the February 1941 issue of *Astonishing Science Fiction* magazine. The story, aptly called “And He Built a Crooked House,” concerns the young architect Quintus Teal. In discussion with a friend, Teal questions, “Is the static geometry of Euclid the only mathematics?” and then realizes that he can save on land charges by building a house in the form of a tesseract – which is the name coined by British mathematician Charles Hinton for the three-dimension projection of a four-dimensional hypercube.

A tesseract looks rather like a cross made up of eight cubes, with the central column being composed of four cubes stacked on top of each other and with the third cube in the column having four cubes attached to each of its outer faces (Fig. 1.2). Having built a house according to the tesseract design an earthquake occurs in Heinlein’s story, and this literally shakes and folds the three-dimensional tesseract shadow-house into an actual four-dimensional hypercube. Upon entering the house Teal and friends become trapped but are eventually able to escape before another earthquake strikes and the house completely disappears – “It simply fell through into another section of space,” concluded Teal.

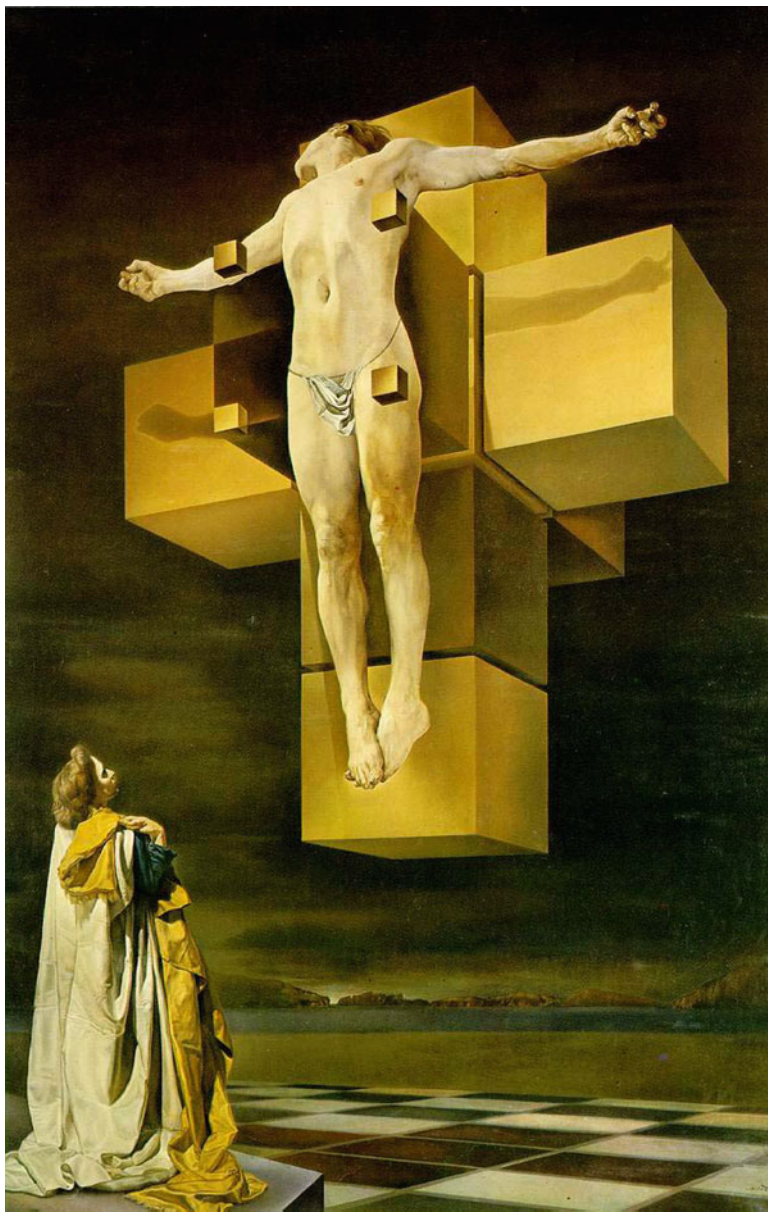


Fig. 1.2 One of the most outstanding renditions of a tesseract is that by surrealist artist Salvador Dali (1904–1989). In his 1954 painting “The Crucifixion” (*Corpus Hypercubicus*) Dali depicts the hypercube piercing the Christ figure, suggesting, perhaps, that the historical Christ holds dominion over all dimensions

Perhaps the best-known story to be set in a world having different dimensions to our own is Edwin Abbott's *Flatland: A Romance of Many Dimensions* (first published in 1884). Abbott's novella is on the surface a humorous romp through the geometry of a two-dimensional world – it is inhabited by polygons and lines that live in houses and have a complex set of social norms and conventions. Behind the geometrical doggerel, however, is a scathing satire of Victorian society and religious persecution. One of the key mathematical points that Abbott draws out, however, is that as far as the Flatlander's are concerned there is no such thing as three-dimensional space. Any object actually having three dimensions could only be perceived, by a Flatlander, via the two-dimensional shadow that it casts – they would never be able to see or sense a three-dimensional object.

The LHC and Extra Dimensions

Four-dimensional houses and Flatlands are not constructs that we can experience, at least in our everyday lives, but this is not to say that there are not more, possibly many more, than the three spatial and the one time dimension that we are familiar with. If one could move, somehow, between other dimensions, those outside of our familiar three, then in principle one would be invisible. Such movement has a long science fiction and fantasy writing heritage, but for the practical physicist the first real issue that must be settled is the question of additional dimensions – are there really four, five, ten or more spatial dimensions to our universe? Remarkably, we might be able to decide the answer to this question in the near future. The device that could provide this answer is the Large Hadron Collider (LHC), situated at CERN on the border between France and Switzerland.

The LHC is the world's largest and most powerful atom smasher. Accelerating two oppositely rotating beams of protons around a 27-km ring of super-cooled magnets at close to the speed of light, collisions are engineered within the confines of a series of experimental halls. The task of the collider and the detectors are to create, albeit fleetingly, and then measure the properties of exotic elementary particles. By crashing protons together at incredibly high energies a great spray of atomic debris is generated, but by identifying the kinds of particles present in the debris field the history of the collision can be reconstructed. Indeed, many different kinds of outcomes can occur when two protons collide, and one of these potential outcomes is the generation of a miniature black hole (see the next section). These black holes are so small that they soon evaporate away via what is called Hawking radiation (see Appendix 1), and as they decay they will generate a veritable cornucopia of exotic particles that will, in the words of one CERN researcher, make the detectors light up like a Christmas tree. The decay of a miniature black hole can hardly be missed, and remarkably the detection of any such decays will indicate that there must be more than three spatial dimensions to our universe.

This incredible possibility of revealing extra dimensions comes about because of an energy constraint. While the LHC is the most powerful collider in the world,

it will only be capable of generating microscopic black holes if we inhabit a universe in which there are numerous extra dimensions. This extra dimension requirement, however, is one of the key ideas behind what is known as string theory – a long-running but still speculative theory that considers elementary atomic particles to be vibrating loops rather than billiard ball-like structures.

There are many different versions of string theory, but they all require that microscopic extra dimensions exist, and anywhere from 10 to 11 and even 22 extra dimensions have been postulated by theorists. The additional dimensions required in all versions of string theory are not visible to us since they are effectively curled up into microscopic domains. Indeed, the extra dimensions are reduced to the so-called Planck scale of 10^{-35} m, at which point quantum gravity effects must dominate any physical interactions. For the LHC researchers, however, the more microstructure dimensions there are to space the better, since the energy threshold at which miniature black holes might be created decreases the higher the number of dimensions there are.

Although it might be possible in the near future to say that additional spatial dimensions, beyond the familiar three, must exist, it is altogether another thing to try and physically move between such dimensions. Possibly a Type II impossibility, dimension shift as a means of achieving invisibility is firmly fixed in the science fiction/fantasy realm at the present time – and indeed, will likely remain so for the foreseeable future. This being said, futurist, writer and professor emeritus of mathematics at San José State University, California, Ruby Rucker has, in his own words, “recklessly” predicted that a way will eventually be found to change the very topology of space (see Chap. 5), gaining access to the variously curled-up dimensions predicted to exist by string theory. Lest he be charged with the transgression of making idealistic speculations, Rucker has suggested that such technologies might become available, “before the Year Million.”

Irreversibly Invisible: Black Holes

Question: When is a door not a door?

Answer: When it is an event horizon.

In the list of most extreme, bizarre, and least understood objects in the universe, black holes must be somewhere near the top. They are objects that on the one hand are simple enough to be described mathematically but on the other have attributes and interior regions where we know our current understanding of physics must fail.

Black holes are all about the extremes of gravity. Whenever enough matter is squeezed into a small enough space a black hole will form, and at the same time set up a one-way trapdoor about itself. The trap door is called the event horizon, and once you have crossed its threshold there is no way back – the black hole has you forever. Indeed, a black hole is the ultimate invisibility cloak, hiding its content

from view to any observer in the entire universe. There is the drawback, however, in that this invisibility is irreversible. Any message (or secret code or person or spaceship) that is sent into a black hole for safe keeping is gone forever and cannot be retrieved – not even by the original sender.

In spite of its name, upon crossing an event horizon absolutely nothing will physically happen to you – nothing, that is, until you try to pull away from the certain death that lurks in the twisted gravitational hell of the central singularity. The so-called escape speed that is required to move away from the black hole, without falling back in again, is equal to the speed of light at the event horizon, and this is why crossing it will be a one-way trip, since no corporal object can travel at light speed or faster.

What happens at the very center of a black hole, the singularity, is a profound mystery. All that can be currently said is that known physics must break down there since, at least currently, there is no agreed upon theory of quantum gravity. Indeed, an agreed upon theory of quantum gravity is the Holy Grail of modern physics. And if found, and not everyone thinks it ever will be, it will tell us how gravity behaves on the minuscule scale of the quantum and subatomic world.

How long is forever? Here's a question for the philosophers to ponder and the guardians of deep secrets to worry over. A massive black hole will certainly protect its internal secrets from outside recognition for an extraordinarily long time, but in a remarkable series of research papers published during the 1970s Stephen Hawking (Cambridge University, England) showed that black holes are not entirely black. Indeed, Hawking found that a steady emission of radiation, appropriately called Hawking radiation (see Appendix 1), will issue from the event horizon region, and this will cause a given black hole to slowly evaporate. It might take an incredibly long time for a black hole to evaporate, but the key point is that they are not eternal, and so the question arises as to whether they truly keep their internal secrets safe forever. The present consensus is that any cryptic message hidden within a black hole can, in principle, be read – although the script will be extremely difficult to interpret, making it a Type II impossible problem rather than a Type III unsolvable one. Essentially, the argument goes, the internal secrets of a black hole are encoded, somehow, within the Hawking radiation, and in principle, one could monitor these diaphanous emissions with the aim of deciphering the code.

Although most physicists accept that Hawking radiation should exist, it has never actually been detected (see Appendix 1, however), so for the moment black hole secrets are truly safe from prying eyes. Looking to the deep future, however, Stephen Hsu (University of Oregon) has considered the possibility of creating eternal black holes that do not emit Hawking radiation. Posting his ideas on the arXiv pre-print webpage for July 17, 2010, Hsu draws upon the idea that black holes are the time-reversed companions of white holes. That is, if one runs time backwards a black hole will appear to be a matter-generating region, spewing material into space rather than locking it away. Such an object would be a white hole. Astronomical observations indicate that while black holes exist, there is no indication that any white holes have formed within the observable universe, which is not to say that white holes don't exist. Be all this as it may. The result of Hsu's investigation was the realization

that if one could find a star that was about to collapse into a black hole, and then, at exactly the right moment fire into it a precisely tuned burst of radiation, the subsequent Hawking emission process could be quenched – an eternal, none-radiating, black hole would thus have been formed and the ultimate safe haven for never-to-be-revealed secrets created.

Time Shift

Rather than shifting (somehow) into other spatial dimensions to achieve invisibility one might alternatively try to shift in time instead. The key idea here is to avoid being at the same spatial coordinates, at the same time, as the observer one is trying to hide from. It is the ultimate, “I was there, but when were you” trick.

The quintessential story concerning time travel is, of course, that of *The Time Machine* by H. G. Wells. Written in 1895, *The Time Machine* was penned 8 years after *The Invisible Man* and considers possible futures for humanity.¹ The time traveler also disappears at the end of the novel, choosing to live in a time different to that of his contemporaries. In this case invisibility is achieved because while located in the same three-dimensional space as everyone else (indeed, it is the same Earth), the time traveler is not there at the same time as everyone else.

Is time shifting possible? The answer to this is probably no, but the full situation is far from settled. Stephen Hawking observed a number of years ago that the absence of tourists from the future is a good argument against backwards time travel. This argument holds true, of course, until we actually meet a time traveler in our future. This being said, it is probably safe to say that most physicists hold that time travel is not possible, although there are some researchers who have found exotic mathematical schemes within the theory of general relativity that seemingly open the doors to a form of time travel – on this topic only the future will tell us what can be achieved.

In the Eye of the Beholder

Writing in 1879, the German astrophysicists Johann Zollner (1834–1882) wrote a remarkable, albeit reputation-killing book called *Transcendental Physics*. The text was concerned with the results from a series of experiments Zollner had performed with the American medium Henry Slade. Indeed, Zollner presented his book as a

¹It is remarkably to think that between 1887 and 1901, Wells pretty much explored and/or defined all of the major themes of present-day science-fiction writing: from alien invasion in *War of the Worlds* (1898) to space travel in *The First Men in the Moon* (1901) and the misguided genetic generation of monsters in the *Island of Doctor Moreau* (1896).

physical investigation into the paranormal, including clairvoyance, matter transfer, ghosts and natural spirits. The main thrust of Zollner's argument was that all these kinds of mystic phenomena could be explained if ghosts (spirits of the dead) inhabited a "real external world" composed of four spatial dimensions. On this basis, he argued that what we see as ghosts are really just shadows projected into our three-dimensional spatial realm.

Well, as we just mentioned, the book didn't do Zollner's career any good, and his subject, Henry Slade, was a controversial figure in his own right, being charged on a least one occasion with fraud. In spite of Zollner's claims that experimental controls were set in place he was possibly the first (but unlikely the last) physicist to be fooled by a showman's slight of hand – the media savvy trickster Uri Geller would be a more recent example of a performer who can apparently achieve feats of psychic metal-bending and mind-reading acts under supposedly controlled experimental conditions.

The ultimate in the eye-of-the-beholder genre of stories is probably that of the *Emperor's New Clothes* by Danish author and poet Hans Christian Anderson. Written in 1837, this well known tale concerns the vanity of conceit. Under the ruse that he is using a special magical cloth that is invisible to any stupid or inadequate person, a tailor convinces the naked king that he is, in fact, wearing a magisterial gown. The king and his courtly retinue, although clearly seeing the nakedness, refuse to admit it since it would publicly reveal their mental ineptitude. It ultimately takes an innocent child to reveal the absurd truth of the situation. Here invisibility is used in a raw emotional sense, not so much for what it hides but for what it reveals.

Absurd invisibility is invoked by British author C. S. Lewis in his novel *The Voyage of the Dawn Treader*. Written in 1952 this story once again sees the Pevensie children in the mythical realm of Narnia, and at one stage they encounter the invisible Dufflepuds. The Duffers, we are told, were once dwarf-like creatures who, out of sheer malice, had been transformed by the magician Coriakin into monopods. Unhappy with their transformed appearance, having just one large leg and an enormous foot, the Duffers asked Coriakin to make them invisible, in order to hide their ugliness – and Coriakin was apparently only too happy to oblige them. The Pevensie children eventually restore the Duffers to visibility (and then rename them Dufflepuds). In this context invisibility is being invoked as a form of emotional cloak; the Duffers feel so disfigured that they would rather not be seen by anyone; this is invisibility as the antithesis of Anderson's naked emperor.

The Dufflepuds described by Lewis have an actual basis in medieval and ancient folklore, where it was suggested that monopods really did exist. Indeed, folklore is rich with mystical beings that have the ability to control their appearance and shape, and the variously named fairies, goblins, banshees, elves, pixies, and many more, can all apparently invoke invisibility, disappearing from our world into their own spiritual realm. To our ancestors such beings were very much alive, and in the late sixteenth century "Fairy Trials" were commonly held on the Island of Sicily. For those charged with the crime of communicating with fairies a sentence of exile or jail could be imposed.

The Cottingley Fairies

Even in modern times, fairy stories can still invoke newspaper headlines and claims of the miraculous. One of the most famous such events is that of the Cottingley fairies, although perhaps the most amazing part of this story is that it became a story at all. To the modern, skeptical, and more jaundiced eye the fairy images are so clearly contrived that they couldn't possibly be anything but fakes (Fig. 1.3). The images of supposed dancing faeries were taken by two young girls, Elsie Wright (aged 16) and Frances Griffiths (aged 10) in Cottingley, England, in 1916. The pictures were taken with a standard film camera and would probably never have entered the public domain if they had not been brought to the attention of Sir Arthur Conan Doyle.

Doyle was a well-known evangelical spiritualist. Like Johann Zollner before him Doyle advanced the idea that the spirits of the dead inhabited some extra-dimensional realm, and he also advanced the idea that those entities such as fairies, goblins, and elves were able to directly cross a psychic dividing line, allowing them to step over into our consciousness, and to take-on "a certain degree of materialism." Indeed, in 1922 Doyle published a short book called *The Coming of the Fairies*, in which he presented a series of letters, newspaper clippings, and thoughts on natural spirits and the Cottingley fairy pictures. Although Doyle appears to believe that the



Fig. 1.3 One of the Cottingley faerie pictures. This picture shows 10-year-old Francis Griffiths and 'friends' taken in 1916

photographs were not fake, he does write in the preface that the “narrative is not a special plea for that authenticity.” The ultimate twist to the Cottingley story took place in 1981, some 51 years after Doyle’s death, when Elsie Wright (then 80 years old) finally admitted to a documentary maker that the images had not so much been faked but were simply photographs taken of themselves with colored card and cloth cutouts of faeries arranged about them.

We learn many things about human nature from the Cottingley saga. Firstly the images remind us that what we think we might see, whether directly by eye or in a photograph, is not necessarily real. Second, we can be fooled by photographs: the Cottingley negatives were actually studied by several photographic experts (even one from Kodak) and not one of them felt that they had been faked. In their defense it should be added that technically the negatives hadn’t been doctored; it was the scene and the circumstances under which they had been exposed that were faked. Thirdly, people often see what they want to see. Doyle, for example, accepted the authenticity of the images since he wanted to find proof for his spiritualist beliefs. The moral that we should take away from this story is that seeing is not just about the process of forming an image on our retinas; it is also about our psyche, our expectations, our hopes, and often our innermost desires. Advertising agencies have never been shy in exploiting these all too human aspects of seeing, and mind manipulation is, in fact, the name of their game.

The Multiverse

Although the notion that fairies and ghosts are somehow entities crossing over from some other world or universe is, well, a fairy tale, since the idea has no scientific rationale to it, this does not mean that physicists haven’t embraced the idea of multiple universes. Indeed, our observable universe is generally taken to be just one very small segment of a vast assemblage (possibly an infinite assemblage) of universes. Each one of these so-called parallel universes can in principle differ from our own with respect to the allowed physical laws and physical constants.

Depending upon the exact version of the multiverse being invoked (and there are many different versions), some physicists suggest that all conceivable laws and physical constants are represented somewhere in the multiverse. This form of thinking can be used to explain why some aspects of our universe seem so finely (and thereby improbably) tuned to allow life to have evolved (Fig. 1.4). Under the guise of the Anthropic Cosmological Principle this version of the multiverse posits a form of the Goldilocks effect at play. Here the idea is that in some universes the laws of physics are just right for life to emerge, while in others the laws essentially conspire to act against life ever evolving. Our existence, therefore, is essentially a selection effect and the universe is the way it is because we wouldn’t exist to see it if it wasn’t.

A different version to the multiverse idea is that proposed in 1957 by Hugh Everett III in the so-called many-worlds interpretation of quantum mechanics – a view

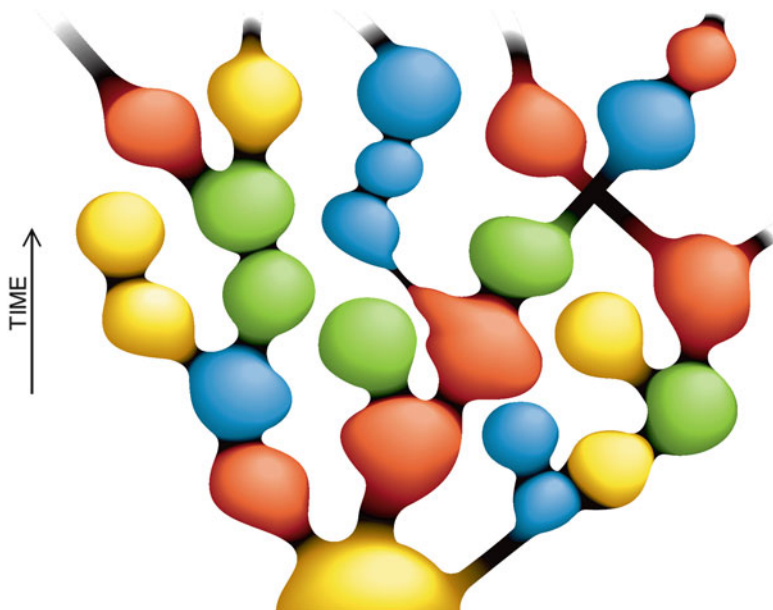


Fig. 1.4 A schematic view of the expanding multiverse as envisioned by Stanford University physicist Andrei Linde. Here the idea is that the universe continuously divides into more and more sub-universes, each of which expands and grows over time and in which the constants of physics, and even the number of dimensions, can be different from our (very) specific visible universe (Image courtesy of Andrei Linde)

more recently and much more vociferously advocated by Oxford University physicists and quantum computation theorist David Deutsch. The multiverse invoked by Everett assumes that all universes are structurally identical (that is, obey the same laws of physics and have the same values of physical constants), and all share a common time parameter. The various universes, however, can and indeed will be in very different states.

Everett's ideas concerning the many-world interpretation grew out of the physics of quantum mechanics. In the quantum realm certain observations cannot be predicted absolutely – only the probability of an outcome can be enumerated. Under these conditions, according to Everett, all of the various possible outcomes do actually come about at the same time, but each outcome is realized in a different universe. This provides the comforting thought that in at least one, or maybe even a few, universes the writer of this book is a wealthy best-selling author. The other point, of course, is that if one could slip into another universe, in a process similar in idea to the dimension shifting described earlier, you would then be completely invisible to all observers in the universe that you shifted from. As with spatial dimension shifting, however, the way in which one might travel between multiverses is entirely unclear and may also be a Type III impossibility.

The Trojan Horse

Invisibility can be achieved in innumerable ways, and it is often said that the best way to hide something, that is, make it invisible to our senses, is to put it in plain view. Animals exploit this means of cloaking through the use of camouflage and shape mimickery. A slight twist on this strategy is to place the object to be hidden inside a harmless looking body that has been left in plain view. This is the classic Trojan horse ruse, and also, perhaps, the ultimate in the “Treachery of Images” as depicted by the painter Magritte (Fig. 1).

The Trojan War was set in a time that was ancient history even to the ancient Greeks and supposedly took place sometime in the thirteenth to twelfth century B.C. The renowned bard Virgil recounts the events of the war’s closing stages in his epic poem the *Aeneid*. Having failed to break into the city of Troy by force, the Greek army pretends to abandon the battlefield and begins to sail away. In the meantime a large wooden horse has been constructed in which are hidden some 30 troops and two spies (hidden in its mouth). On being duped into thinking that the siege had been abandoned the Trojans pull the oversized wooden horse into the city, thinking it to be a triumphant battle trophy. Well, of course, later that night the soldiers creep from the horse, open the city gates and the returned Greek army enters to destroy its Trojan enemy.

In the modern era, the Trojan horse approach of invisibility is the favored mode of attack by the malicious computer hacker, the bane of the internet, who is intent upon gaining access to a remote target computer. The idea is simplicity itself, with the malware (malicious software) being embedded and hidden (and therefore invisible) in an openly flaunted, harmless looking piece of computer code. Once downloaded onto a target computer, however, the malware goes to work and corrupts the host system. Access to an otherwise secure area of the computer has been gained by hiding the virus within a harmless looking shell. It is a trick, but it is a trick that works very easily, and the victim essentially does all the hard work for the hacker.

The Vanishing Pyrex Block

How can a Pyrex glass block be made to disappear before your very eyes, while also sitting in plain sight? Sounds impossible? Well not so, as the following liquid Trojan horse experiment will reveal. All that is needed is a kitchen, or at least a table, a glass beaker, a Pyrex block, and a bottle of vegetable oil. The trickiest part of the experiment is actually finding a Pyrex block – anything made of Pyrex will actually do – a plate, or a small cup; for convenience, however, we will keep using the term block for the Pyrex item to be made invisible.

Place the Pyrex block in the glass container. At this stage the block is certainly visible. The subterfuge is completed, however, by adding the vegetable oil. Fill the glass container with the oil and you will see that as it covers the Pyrex block it vanishes from view – invisibility in plain sight achieved. The container full of oil is a Trojan horse for Pyrex blocks.

What is going on in this experiment? Obviously the vegetable oil is the key ingredient, but what does it do? The oil actually suppresses two effects: the reflection of light from the surface of the Pyrex block and the refraction of light through it. These effects, reflection and refraction, are exactly the two processes that allow us to see anything solid. The physics of these phenomena will be described in Chap. 2, but for now the point is that by covering the Pyrex block with oil reflection from its surface is suppressed, and by using vegetable oil the Pyrex block has been immersed in a fluid with an identical refractive index (explained later). What this means is that light will pass through the fluid and the block at the same speed and consequently there will be no refraction, or bending, of the light as it crosses the oil-to-Pyrex block surface boundary. By neither reflecting nor refracting light at its surface the Pyrex block will be indistinguishable from the surrounding oil and accordingly it will be invisible to our eye.

Wood and Strutt Explain

American physicist Robert Wood was a no-nonsense kind of man. His area of expertise was optics, and his passion was to debunk hoaxers and hoaxes. It is perhaps a little surprising, therefore, given his interest in exposing pseudoscience that Wood appears to have written the first scientific paper on invisibility. The article appeared in the August 1902 issue of the journal *Physical Review* and is just two pages long. Entitled “The Invisibility of Transparent Objects,” the article begins by explaining the vanishing trick just described. Wood writes, “A transparent body, no matter what its shape, disappears when immersed in a medium of the same refractive index.” To this he adds, “Could a transparent solid substance be found, whose refractive index and dispersion were the same as those of air, it would be absolutely invisible.”²

In his research paper Wood makes reference to an 1884 article written in the *Encyclopedia Britannica* by British physicist John Strutt (perhaps better known through his honorary title of Lord Rayleigh). In a brief footnote in his review, Strutt explains that “perfectly transparent objects are only visible in virtue of non-uniform illumination,” and to this he adds that “in uniform illumination they would be absolutely invisible.” In this sense, Strutt is saying, we only see uniformly transparent objects, such as glass blocks, because they are invariably viewed under conditions where the light is coming from one very specific direction or another – and rarely, if ever at all, are they seen evenly illuminated from all directions. This latter condition does raise the philosophical question, “How do we know when we haven’t seen something?” The answer, of course, is “When we fall over it.” The conditions for obtaining invisibility, Strutt suggests, however, might be realized on top of a tall

²This, recall, was exactly the idea presented by H. G. Wells in the story of *The Invisible Man*.

monument in a dense fog – from which one generates the fanciful vision of shadowy characters attempting to scale Nelson’s Column in Piccadilly Circus, London, during a classic, Victorian-era pea-souper smog.

Rather than go to such extremes, however, Wood, ever the experimental physicist, set out to make a uniform illumination chamber in his laboratory. To this end he took two glass bowls and painted their interior surfaces with “Balmain’s luminous paint.” He then made a small viewing hole in the base of one of the bowls and joined the two dishes together to make a sphere. Before joining the two bowls together, however, he first exposed the luminous paint to strong sunlight and placed a small glass sphere within its interior. With the two bowls joined together, Wood, now ensconced in a darkened room, looked through the small viewing aperture. He writes, “A uniform blue glow fills the interior of the ball [the two bowls] and only the most careful scrutiny reveals the presence of a solid object [the glass sphere] within it.” Invisibility, well very nearly so, had been achieved in the laboratory for the very first time – a transparent object was rendered un-seeable.

Blending In

Contemporary Japanese artist Yayoi Kusama has exploited the idea of visually blending into the background, making her body or parts of her body invisible, in many of her works (Fig. 1.5). Here the idea of invisibility is simulated by dressing the subject in the same color and with the same design as the background. The inherent problem with this approach to invisibility, however, is that as soon as the subject moves the invisibility effect is completely lost.

The technique employed by Kusama is actually very versatile and commonly used on film sets and in television studios. We are all probably familiar with the view of our local TV weather forecaster standing, apparently, in front of a map of the approaching weather systems. In reality, of course, the person is standing in front of a uniformly colored blue screen, and the map is added digitally over the background. Provided the forecaster isn’t actually wearing any blue clothing, then they will appear in front of the background map. With the blue screen technique the idea is for the actor or presenter to be seen. The full invisibility effect, however, can be achieved and maintained, no matter how the person moves, if they are dressed in a full bodysuit of uniform color in front of the same uniform color background. Many theater companies have successfully exploited this effect and used actors in black bodysuits working against a black background to manipulate, in a seemingly invisible manner to the audience, colored objects on stage.

The idea of camouflage and to a certain extent an adaptation of the Trojan horse ruse is that adopted by the wolf in the Grimm’s fairy tale *Little Red Riding Hood*. Here the wolf attempts to disguise himself as Riding Hood’s grandmother, her friendly exterior (actually her clothing) acting as a cover for the wolf and his deadly intentions. Cloaking in the Red Riding Hood style, by which an object is camouflaged through the projection of a benign object onto its surface, will be described



Fig. 1.5 The body of Japanese artist Yayoi Kusama blends into invisibility. Here the foreground and background are made indistinguishable by the careful use of color and pattern selection

further in Chap. 3, but, for the moment, we simply note that the tactical advantage of making, say, a deadly tank look like a harmless milk truck is fairly obvious. Some insects as well as a number of harmless snakes have successfully adapted to an inverse Red Riding Hood camouflage ruse – natural selection acting to make them look the same as a poisonous member of their species. In essence this strategy is the equivalent of making a harmless milk truck look like a deadly tank. The advantage, of course, is that by looking like their poisonous relatives the harmless insect/snake/animal is left alone by its potential predators. Other insects openly flaunt their poisonous nature by adopting highly colored and visually dramatic outer skins – the easier to be recognized and remembered by. Sometimes just being seen is enough to be left well alone – invisibility (of sorts) achieved by being overtly obvious, or simply by looking too deadly a target to engage.

The desire to make military equipment invisible to the enemy, whether in space or on land, has resulted in one of the great conspiracy myths arising out of World

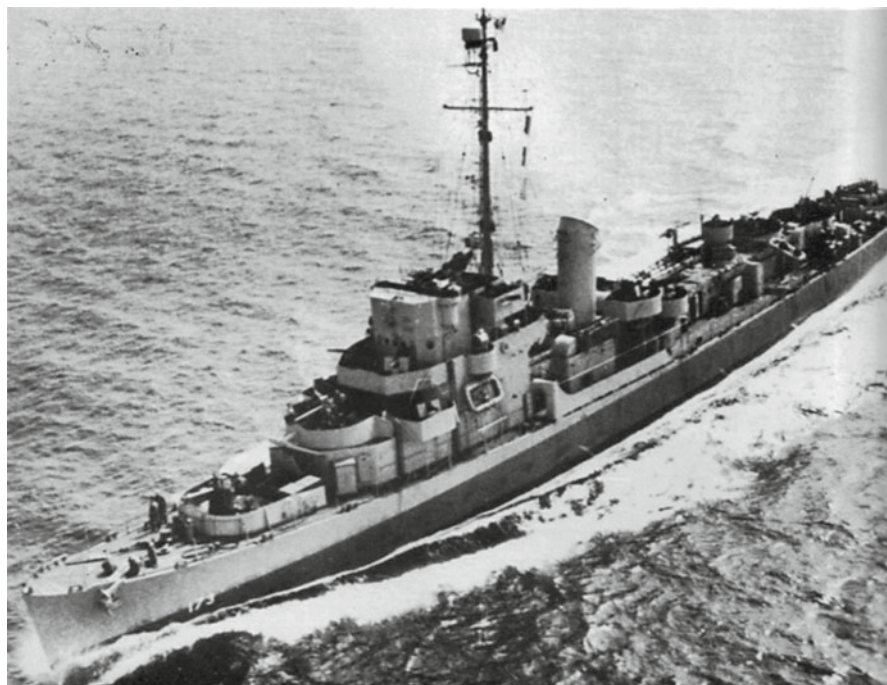


Fig. 1.6 The *U. S. S. Eldridge*, which was supposedly rendered temporarily invisible and teleported through space in 1943

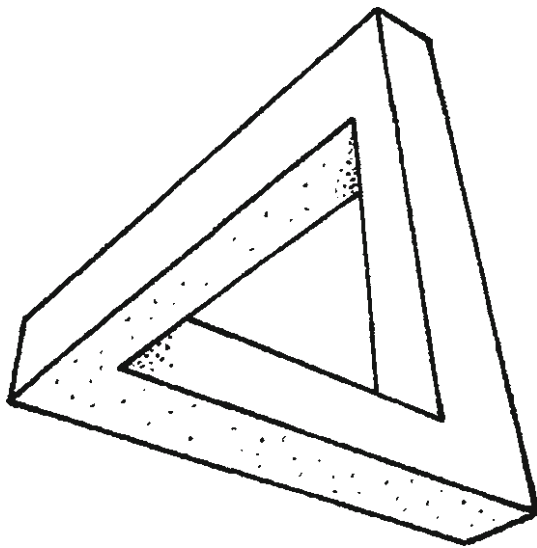
War II. Indeed, the Philadelphia Experiment is one of those urban legends that, in spite of its inherent ridiculousness, just won't go away.

The story dates back to 1943 and involved the apparent disappearance of a naval ship from a Philadelphia shipyard in Pennsylvania. It must be said that this story has been debunked many times, claims that the *U. S. S. Eldridge*, a destroyer escort ship (Fig. 1.6) was rendered temporarily invisible (and some versions of the story claim that it was also briefly teleported to Norfolk, Virginia). What was most likely being conducted on the *Eldridge* was a degaussing program by which the ship was being rendered 'invisible' to magnetic mines. Whether this is the explanation or not, for the *Eldridge* saga it is a reminder of the fact that being invisible is an attribute that operates on many different levels and need not apply to visible light alone – in a sea full of magnetic mines, being magnetically invisible has a much greater survival advantage than being optically invisible.

Seeing the Impossible

Not only can the brain be fooled by tricks of the light, camouflage and Trojan horses, it can also be fooled by geometry, and remarkably we can actually *see* impossible things. Swedish artist Oscar Reutersvard (1915–2002) used impossible figure

Fig. 1.7 The impossible triangle of Lionel and Roger Penrose



designs in many of his works, beginning in the 1930s, but it wasn't until some two to three decades later that the genre really became popular. Figure 1.7 shows an image of the impossible triangle first described by Lionel and Roger Penrose in 1956. The shape is an impossible three-dimensional object (that is, it cannot be physically made in three-dimensional space), but as a drawing on a two-dimensional surface it displays an odd, even unsettling, combination of local possibility (each segment and corner of the triangle is correctly drawn) coupled to global impossibility. The eye and brain continuously reassess the image, not being able to find a stable interpretation. Dutch born graphic artist M. C. Escher (1898–1972) is especially well known for developing and using variants of the Penrose impossible figure in his works, depicting ill-omened walkers apparently trapped upon Sisyphus-rising treadmills and endless staircases. Such figures now routinely adorn many a student dormitory wall.

How we see the world depends upon many factors, but sometimes what we see is not what is actually displayed. Spanish-born painter Salvador Dali made great use of such illusions in his surrealist works, and perhaps one of his finest pieces is the 1940 oil on canvas, “Slave Market with the Disappearing Bust of Voltaire.” In this painting the bust of Voltaire is encoded within the profile and shadow play of two female figures. The brain struggles to find stability within the presented image and alternately switches between seeing the hollow-eyed face of Voltaire and the two scurrying women. Dali said of this painting that he was attempting “to make the abnormal look normal and the normal look abnormal.” He succeeded!

Optical illusions are both fun and revealing – yet psychologists are still unsure why they work and what it is they tell us about the brain. But work they do, and German astrophysicist Johann Zollner, whom we met earlier with respect to his book *Transcendental Physics*, discovered a most remarkable illusion in 1860 (Fig. 1.8).

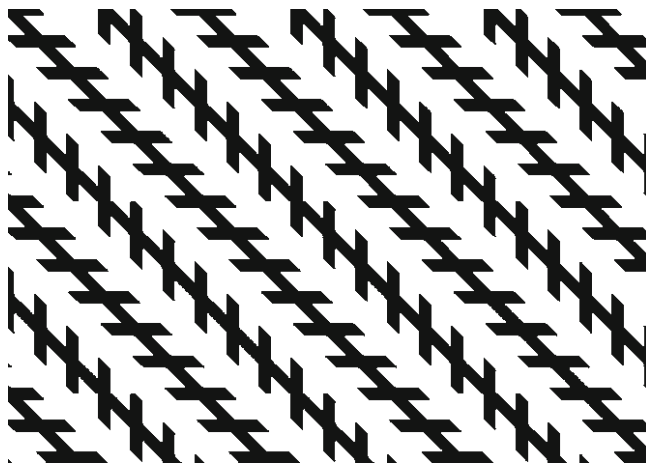


Fig. 1.8 Zollner's illusion. The short horizontal and vertical bars make the longer, parallel diagonal lines appear to splay apart

Indeed, optical illusions trick the brain into seeing something that isn't actually drawn, and in the Zollner illusion it is the addition of the short horizontal and vertical lines that make parallel diagonal lines appear to diverge. Interestingly and indeed mysteriously it turns out that the illusion can be destroyed under some line and background color combinations.

Impossible figures are not quite the same as optical illusions, and it is again interesting that we can see such figures at all. The human brain is an incredible evolutionary adaptation, and it continually strives to find order within the chaos of light and motion that surrounds it; and yet, the brain can be fooled. Indeed, the physical workings and methodologies of the brain, we should remember, are still a profound physiological mystery.

As far as the author is aware there are no examples of an impossible figure being used to hide, that is, make invisible, an otherwise recognizable object – a topic seemingly ripe, therefore, for further exploration. The closest approach to this latter ideal is perhaps that of the anamorphic print. Anamorphic (derived from the Greek for *re-transformed*) pictures are constructed according to some exact projection that requires the observer to adopt a very specific viewing angle and/or use a specially curved mirror to reconstruct the image. "The Ambassadors" painting, produced circa 1485 by Hans Holbein the Younger, contains the anamorphic projection of a human skull across its lower center. To view the skull correctly the observer's viewpoint must be nearly flush with the canvas and looking either up from the lower left, or down from the upper right. Hungarian-born painter Istvan Orosz (born 1951) has also explored anamorphic projections in a number of his works (Fig. 1.9) and has commented, "There are things I can imagine and draw. There are things I can imagine but cannot draw. But, could I draw something that I cannot imagine? That interests me greatly."

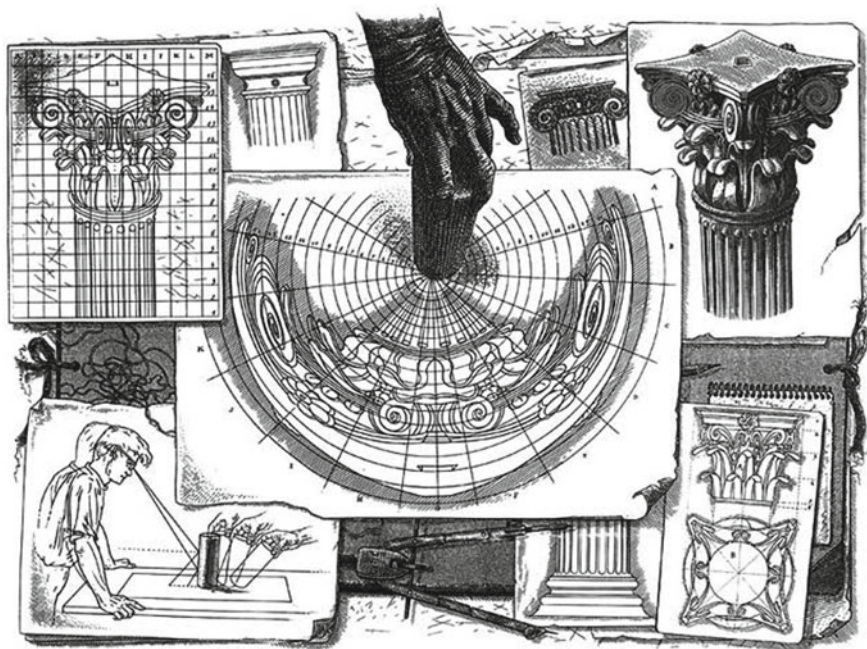


Fig. 1.9 Design outline for the print “Mirror Anamorphosis with Column,” by Istvan Orosz. In this anamorphic design the column is resolved when the print is viewed via reflection in a cylindrical mirror

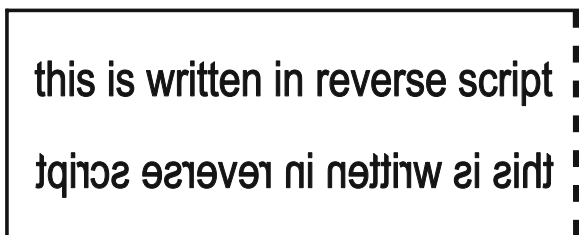


Fig. 1.10 *Top*: an example of the reverse mirror-writing script used by Leonardo da Vinci. The lower line was generated by taking the first, ordinary line of text, and performing upon it a horizontal and then a vertical flip rotation. To read the mirror-written line correctly, place a flat mirror at right angles to the page along the dashed line

The distortion in an anamorphic figure need not be highly pronounced to cause non-recognition, and the great Italian renaissance painter and inventor Leonardo da Vinci (1452–1519) knew this well. Accordingly, da Vinci wrote his personal notes and journals in a mirror-writing script than ran from right to left across the page (Fig. 1.10). In this manner, the text reads correctly only when seen reflected within the surface of a flat mirror. The effect seems charmingly simple, but it is enough to

thoroughly confuse the eye and brain when attempting to read the script without a mirror. In essence, to ordinary sight the mirror script is devoid of meaning and the message is made invisible.

The Ultimate Optical Illusion

The human eye is a marvel of natural selection and evolution. For all its remarkable adaptations, however, it is not constructed in a manner that makes obvious or logical sense. Indeed, as Richard Dawkins has beautifully explained in his (highly recommended) book *The Blind Watchmaker*, the human eye is wired up backwards. The light-sensitive parts of the eye's retina, the cones and rods, actually face away from the lens. It is as if the director on a movie set is trying to film a scene by using a large mirror and having the camera record the reflected performance of the actors.

In practical terms this arrangement means that the neural pathway connecting the brain to the eye's light receptors physically obscures and blocks from view some of the image produced by the lens. The various pathways all come together to make the optic nerve, and where this bundle of fibers passes through the retina there is a blind spot – a small zone that has no light receptors at all. The remarkable thing about our vision is that we do not normally see the blind spot. The brain literally fills in the visual gap it produces by using the other eye's image, as well as the colors and patterns surrounding the optic nerve as its guide.

The blind spot can be found if one looks for it in the right way. The first step is to look at Fig. 1.11. Close your left eye and stare at the cross with your right eye only. Now, slowly move the page towards your face. At some point (at about 20 cm from the page) the black dot will magically disappear. In reality, of course, it is at this distance the image of the black dot falls upon the optical nerve (the blind spot)



Fig. 1.11 Test figure for finding the blind spot. Close your left eye and focus on the cross with your right eye. Slowly adjust the distance between your face and the page. At a distance of about 20 cm the black spot should suddenly disappear – this happens when its image falls on the right eye's blind spot

in the right eye. If you continue to move the page closer to your face then the black dot will once again reappear – magic. The phenomenon of the disappearing spot was apparently first described by the French scientist Edme Mariotte in the 1660s. As a remarkable example of convergent evolution, the cephalopod eye, which is almost identical in appearance to the human eye, has its retinal light receptors wired up in the more logical way – facing towards the lens – and accordingly the vistas seen by the octopus and squid have no blind spots for their brains to correct for.

Tripping the Light Fantastic

Is invisibility really possible? Beyond the more straightforward examples of smoke screens, Trojan horses and optical camouflage considered in this chapter, is it really possible to make solid objects, say you or me, a house, a military tank, a battleship or an aircraft disappear from direct view? The answer to this question is an unequivocal yes, but not all of the required technologies are fully in place – yet. How close are we to delivering these technologies? The answer to this question will be explored more fully in the following chapters, but the truly exciting point at this stage is that we are talking of developments being made on a timescale measured in decades rather than centuries. Before exploring such incredible devices, however, we should first consider what is it that they are actually working upon – what is light, how is it quantified, and how does it interact with matter? These topics are to be addressed next.

Chapter 2

A Very Brief History of Light

Invisibility is all about controlling and in some sense fooling light – it is a trick. Not a magic or mystical trick, of course, but a real, scientific slight of hand that must be carried out with great technical skill. If we are going to fool light, however, we must first understand what light is, how it moves through space, and how it interacts with matter. This chapter is all about the story of light and the incredible, many-thousand year history of how humanity has struggled to understand its diaphanous, shifting, and on occasion downright spooky nature.

The story of light builds upon the small and the large, the local and the distant. To understand what light is and how it moves we must literally look to the heavens, and we must also peer into the darkened corners of terrestrial laboratories.

Beginning at the Beginning

“Let there be light!” So begins Genesis, the biblical account of the formation of the universe. Indeed, before there was anything, there was light. This description hardly changes in the modern day, Big-Bang theory of cosmic origins which posits that from the very earliest of moments, through to the first 380,000 years, the universe was a broiling broth of radiation. Light – electromagnetic radiation – bathed all space with one blinding, elemental glow. Indeed, a veritable pea-soup fog of lascivious light pervaded the entire juvenile universe. The cooled remnant of this natal glow is still visible to us today, nearly 13.5 billion years on from the beginning, as the cosmic microwave background radiation.

The universe was born of light, and the galaxies, the stars, the planets, and you and I are its progeny. Incredibly, nearly 10 billion years on from the first moments that signaled the beginning of the beginning, the fecund universe saw the evolution of eyes capable of, well, seeing the light. And then, about 2,500 years ago, humans (the intelligent apes) began to ask what it was that their eyes were seeing.

The Act of Seeing

In Sir Arthur Conan Doyle's *A Study in Scarlet*, Sherlock Holmes chastises his long-suffering companion and biographer, John Watson, for observing but not seeing. Here our sympathies must lie with Dr. Watson, since in reality seeing is easy, while observing is difficult – which was perhaps Holmes's point. Indeed, we do not even need to think about seeing; our eyes and brain do it all for us without conscious effort. What we think we see, that is, observe, is something altogether different.

What is it that happens when we see? This question has a basic answer in that what we see is a representation of the physical world around us. Seeing transmits information: there is a tree over there and a sea lion balancing a ball on its nose next to it. But how has the information from the strange scene just described been obtained? Was the process of seeing, that is, the gathering of the data, a passive or an active one? The ancient Greek philosophers, as one might well expect, were the first to consider this question in detail, and in general two basic ideas prevailed.

The ancient Greek were among the first to encourage active exploration. In this situation it was envisioned that a form of light (an inner fire) was generated within the eyes, and upon exiting the pupils filled the surrounding space with an active 'seeing' agent that would interact with and reflect off of solid surfaces, and then, the reflected light that returned to the eyes would reveal the vista before us – the tree, the sea lion, and the ball. A modern parallel to this idea would be that of, say, radar astronomy, where a specific signal is sent outwards by an Earth-based transmitter to be reflected from the surface of, for example, an asteroid (Fig. 2.1), with the reflected signal then being intercepted by a receiver placed next to the transmitter. In this manner, the radar transmitter is actively providing the receiver with a signal that contains information about the shape, distance and speed of the asteroid. The receiver alone will not 'see' or detect the asteroid – it needs the transmitted signal to be sent out and then reflected back in order to record the fact that the asteroid exists.

If light is not actively generated by the eye in order to illuminate the surroundings, then, so the ancient philosophers argued, perhaps all objects shed light from their surfaces. According to this idea our eyes are passive instruments. Indeed, in this case, our eyes are just receptacles in which the myriad light rays husked off the surfaces of foreground objects are brought together to create an image. We still see the tree, the sea lion, and the ball, but now the image exists independently of any observer actually viewing the scene.

To continue our astronomical observation parallel, the emission idea is similar to that of detecting pulsars. In this case a rapidly spinning, highly magnetized neutron star generates a beam of synchrotron radiation that can only be detected if it chances to be directed towards Earth (Fig. 2.2). Casting a radio beam into space, much like a lighthouse casts a message of warning to sailors on the darkened ocean, so a pulsar's emission is modulated by the rotation of its central star. The information that a pulsar exists is thereby periodically cast off into space as a kind of flickering image. This staccato, pulsed emission process is different from that responsible for our

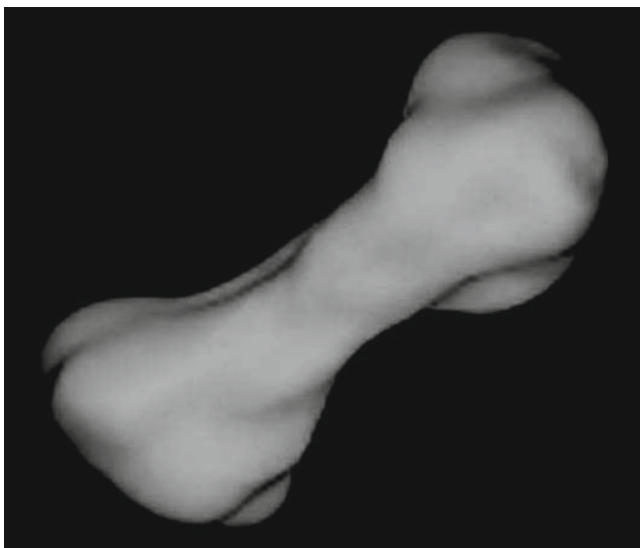


Fig. 2.1 A radar image of asteroid 216 Kleopatra. The image was obtained with data collected by the Goldstone Radio Telescope. The dimensions for Kleopatra are $217 \times 94 \times 81$ km, and the light reflected from its surface reveals that it is predominantly made of a nickel-iron alloy (Image courtesy of NASA)

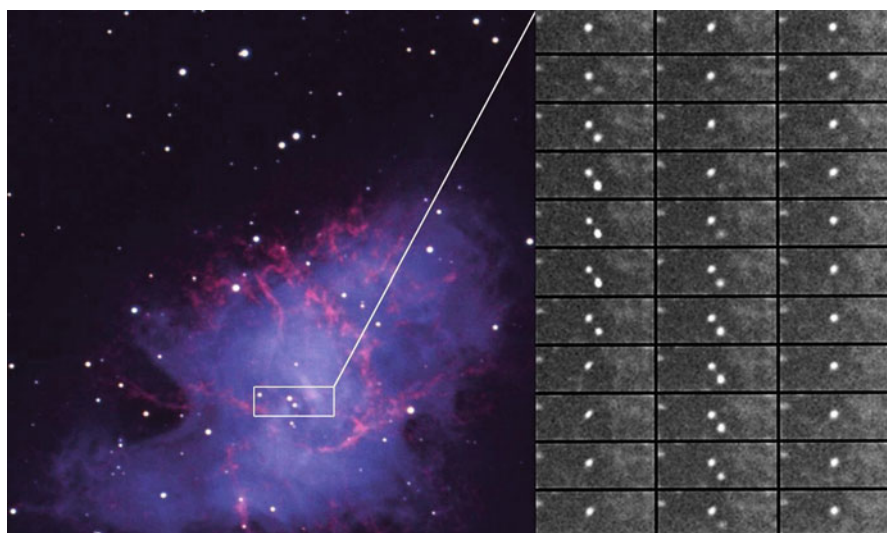


Fig. 2.2 The Crab Nebula and its central pulsar. The on-off sequence of the central pulsar is shown to the right. Whenever the pulsar beam points in the direction of Earth a flash of light is recorded. The host neutron star is spinning rapidly – indeed, it turns on its axis an incredible 30 times per second. By way of comparison, it takes our Sun over 26 days to spin just once on its axis (Image courtesy of NOAO)

image of the Sun, which advertises its optical presence by sending radiation into space equally in all directions continuously.

The philosopher Democritus of Abdera circa 400 B.C. discoursed on the emission idea of seeing, arguing that objects continually cast off images, or simulacra, of their likeness into space. The image, according to Democritus, was impressed upon the surrounding air, and it was thereby transmitted to the eye. Indeed, many of the early Greek philosophers favored the emission idea of vision, but by the turn of the fourth century B.C., the great minds of Plato and Aristotle had begun to shape the notion that a two way process was at work – one might say it was a classic compromise. There was, it was argued, an external image that was continuously husked off the surfaces of objects, but there was also an ‘inner fire’ radiating from the eye that was required to give the image husks substance. It was the observer that literally made the scene come alive. This dual – two kinds of light, internal and external – approach to seeing dominated intellectual thought for the best part of a thousand years, and it was not until the early Middle Ages that alternatives were actively pursued.

Although the classical Greek notion of what light is turned out to be quite wrong, the concept was not damaging to the search for the correct interpretation and development of ideas concerning the workings of light and its interactions with matter. Indeed, the foundations of geometric optics, still valid to this day, were first (at least in part) laid down by the great mathematician Euclid of Alexandria.

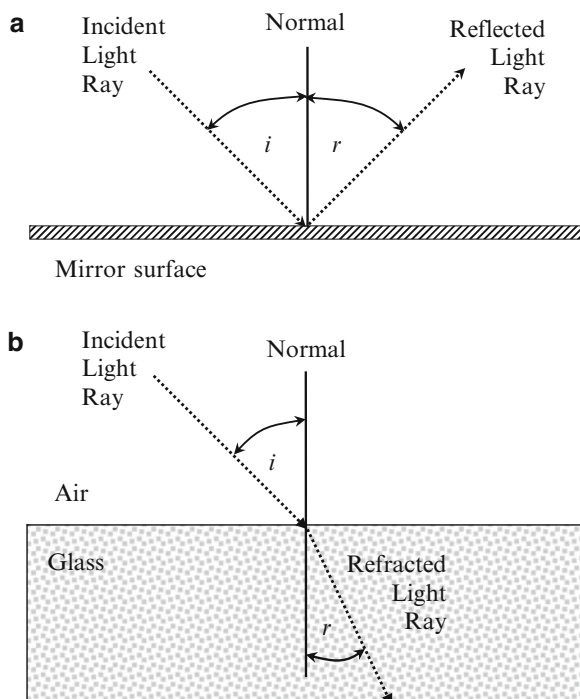
Writing circa 300 B.C. Euclid set out to explain the geometry of vision. Although much of what he wrote need not concern us here, the key, indeed, fundamental, concept that Euclid introduced was the abstract idea of a light ray. Such beams were envisioned to be thin threads of light that traveled in a specific direction (emanating from the eye as far as Euclid was concerned) in a perfectly straight line – this latter attribute being referred to as the rectilinear propagation of light. With this concept in place Euclid then argued that vision – that is, what we see – is composed of those objects situated within a vast cone of light rays, the apex of which is located at the center of the observer’s eye.

The abstract concept of a light ray, and the recognition of the rectilinear propagation of light was a brilliant first step, and although it took another 400 years to work out the details, it was eventually realized that these concepts, if combined with the idea that light rays always travel along the shortest possible path, could explain the laws of reflection and refraction. Hero of Alexandria circa A.D. 100 is often given the nod as the first philosopher to explicitly state the idea of light rays traveling along the shortest path length, and he used this property to prove mathematically the law of reflection, which we will now explain.

The Laws of Reflection and Refraction

The laws of reflection and refraction are fundamental to all optical systems and calculations, and they are cast in terms of what happens to an incident light ray when it encounters a specific surface. They are expressed in terms of the angle of

Fig. 2.3 (a) The law of reflection, with angle of reflection (r) being equal to the angle of incidence (i). (b) The law of refraction. When a light ray passes from a less dense to a more dense optical medium, such as from air to water, the angle of refraction (r) is always smaller than the angle of incidence (i)



incidence, which is the angle that the initial, incoming light ray makes to the surface of reflection or refraction. These laws are:

The law of reflection. The angle of incidence and reflection are equal (see Fig. 2.3a).

The law of refraction. When a light ray passes from a less dense to a more dense optical medium (e.g., from air to water or from air to glass) the angle that the refracted light ray makes with respect to the normal is smaller than the angle of incidence (see Fig. 2.3b).

The law of refraction is clearly more complicated than that of reflection, and indeed we have not as yet explained exactly what we mean by more dense and less dense optical media (but we shall), and neither have we written the law down as a mathematical rule (but we will eventually). In fact, it took many centuries before the law of refraction could be understood and correctly written down as a useful numerical formula.

The Writings of Alhazen

A number of medieval Islamic scholars studied and wrote on optics, but perhaps the most important work was performed by Abu Ali al-Hasan ibn al-Hasan ibn al-Haytham (965–1039), better known to European scholars through the Latinized name of

Alhazen. Alhazen studied copies of the original Greek texts and, more importantly, performed independent experiments. Indeed, Alhazen dissected and studied animal eyes in order to see how they were constructed and to learn how they might function. From such considerations he began to rework the concepts espoused by Euclid, and he specifically argued that the eye was purely a light receptor, and that there was no need for any illuminating rays to exit from within its interior. Light, according to Alhazen, was a purely external phenomenon that illuminated the objects around us, and it is the light reflected from the surfaces of these objects that interacts with the eye to produce an image of the scene.

Although Alhazen reversed the sense of direction of Euclid's light rays, he affirmed the geometric concepts developed by the Greek philosophers. Light rays still travel in straight lines, as indeed we now know they must, and they also travel along the shortest path between the object and the eye. Interestingly, while Alhazen knew that the eye contained a lens and that lenses inverted images, he argued that the lens was not actually used to shape the image in the eye. Rather, he reasoned, the image that we actually see is that produced on the outer surface of the lens. The light didn't pass through the lens to produce an image, Alhazen reasoned, because this would produce an upside down image, and clearly we see the image the right way up. Here we encounter one of the often, even to this very day, forgotten great illusions of seeing – what the eye produces is in fact an inverted image of the world around us, and it is our brain that actually inverts the perceived image. In this way, what we see with our eyes conforms to the orientation of the tactile world we experience.

The Dreams of Dr. Mirabilis

As we move forwards into the twelfth and thirteenth centuries the main centers of learning have shifted to northern Europe, with universities being established by the Church authorities in Paris, Oxford, and Bologna. Although a scholastic approach dominated most of the thinking at this time, we do begin to see the serious consideration of new sciences and new possibilities. One of the foremost writers on optics during the thirteenth century was the Franciscan friar Roger Bacon (c.1214–1294). Given the posthumous title of Doctor Mirabilis, Bacon expounded upon the science of optics in his *Opus Maius*, printed in 1267. This remarkable text, written specifically for Pope Clement IV, covers such topics as mathematics, astronomy, astrology, dynamics, optics and alchemy. Bacon also speculates within the folds of his tome about the possibility of heavier than air flying machines, as well as the construction of new optical devices, the telescope and microscope as we would call them, to improve upon human vision. Indeed, Bacon writes that we may “shape transparent bodies, and arrange them in such a way with respect to our sight and objects of vision, that the rays will be refracted and bent in any direction we may desire, and under any angle we wish we shall see the object near or at a distance. Thus from an incredible distance we might read the smallest of letters and number grains of dust and sand.”

What incredible, imaginative and futuristic speculations these were, and they presage the very same free-thinking that drives current researchers in their quest to construct invisibility devices. Take the known and push the possibilities to the utmost limits of reasoned creativity.

Moving onwards, towards the Renaissance, the various ideas concerning the design and operation of simple optical devices (burning lenses, spectacles and mirrors) begin to converge with our modern-day understandings. What was still entirely unclear then, however, was with what speed did light move. It was known to be exceedingly fast, but how fast, and was it even possibly an instantaneous transmission?

Too Quick, Even for Galileo

To the eagle-eyed, suspicious observer it might have at first seemed that a gang of spies or smugglers had been discovered. For indeed, in the darkness of the Florence countryside two lights were flickering on and off, caught, as it were, in the act of some strange abstract discourse. A conversation of sorts was taking place, but it was the exchange of light rays rather than sound waves that was taking place. Serious scientific investigation was at play. At stake was the determination of the speed of light, and one of the investigators was Galileo himself. The year was 1638.

Galileo turned his thoughts rather late in life to the speed with which light propagates, but he described his nocturnal investigations in his *Discourses and Mathematical Demonstrations Relating to Two New Sciences* (published in 1638). This was one of Galileo's greatest works, and it importantly laid down the foundation ideas to the modern-day theory of dynamics. Within this text, however, Galileo describes how the speed of light might be measured.

The basic idea was that two observers, each equipped with a signal lamp, should position themselves several kilometers apart. The first observer would then open the shutter to his lamp and start counting. As soon as the second observer saw the first light go on, he or she would immediately open the shutter to the lamp. If the two signalers have honed their reaction times to be as fast as humanly possible, then, if the speed of light was finite, the first observer should record a noticeable delay in the illumination of the second observer's signal lamp. Through knowing the distance between the two observers and the delay time for the light signals, the speed of light (at least in principle) could be determined.

Unfortunately, the first experiment to measure light speed was a failure. This was so not because the idea was wrong but because the time delay could not be measured with any degree of accuracy. To determine the speed of light human beings, along with their slow reaction times, would have to be removed from the experiment. It took a further two centuries before this final experimental step was completed.

Hooke's *Micrographia*

Robert Hooke's masterful book *Micrographia* was published by the then newly founded Royal Society of London in 1665. It was, and indeed still is, a beautifully illustrated book, and it stunned an unsuspecting public with its close-up, highly magnified, never-before-seen images of ants, flies, fleas (Fig. 2.4), wasps, gnats, and the minutia of nature. It was a marvel, making the otherwise invisible visible, causing the great diarist Samuel Pepys to eulogize about its content and Hooke to write, "By the means of the telescope, there is nothing so far distant but may be represented to our view; and by the help of microscopes, there is nothing so small as to escape our inquiry." Indeed, with the microscope, as Hooke so clearly realized, human senses could be pushed well beyond their natural limits. The world of the very small and the very distant had been opened up by applied optics. Otherwise hidden and invisible vistas were now accessible to human view.

The publication of Hooke's *Micrographia* marks a turning point in the history of science. Not only did it thrill an awestruck public audience, it also inspired many generations of natural philosophers. A young, moody, and taciturn Isaac Newton read it with enthusiasm in his Cambridge University rooms, although his enthusiasm for Hooke the man was not to last very long.

Within *Micrographia* Hooke described his ideas on light – indeed, he described his ideas on almost everything, causing some readers to complain of its excessively detailed, often obscure and labored text. To Hooke's way of thinking, however, everything in the universe was in motion – nature was never at rest. To this he added

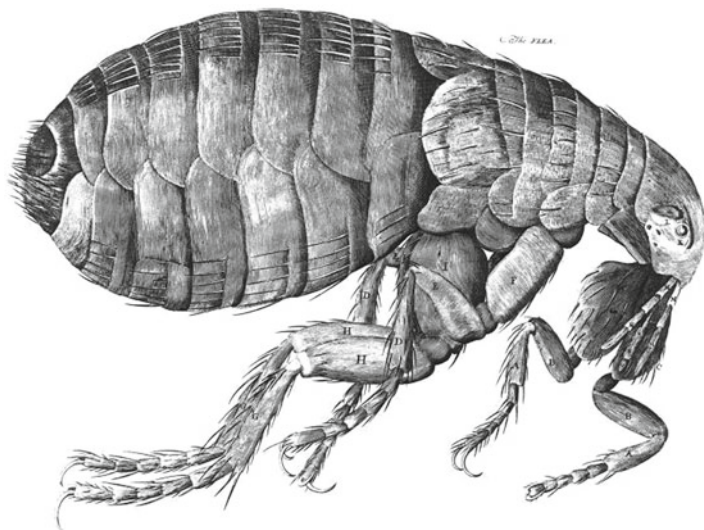


Fig. 2.4 The flea, as revealed by the microscope. Hooke rightly notes that, "The strength and beauty of this small creature, had it no other relation at all to man, would deserve a description." Hooke's original print is a stunning 18 in. (45.7 cm) across

the notion that all motion could be broken down into waves of one form or another; the rising and falling of tides, the phases of the Moon, ripples on a pond, even earthquakes and the repeated passage of the seasons. The world, according to Hooke, hummed like so many vibrations moving along the strings of a complex musical instrument.

Light, in Hooke's view, was an especially symmetric vibration. Unlike the modern-day picture, however, in which light of different colors is viewed as having different frequencies of vibration, Hooke thought that color was a corruption or distortion of white light. Indeed, he imagined light to spread out from an object like so many ripples on the surface of a pond. White light corresponded to those ripples that moved outward along the radii of expanding concentric circles; colors arose under refraction, however, because the path of the ripples were changed and became distorted. In modern terms it could be said that the wave front was no longer transverse to the direction of propagation. The eye, Hooke argued, could pick up on these distortions, and red light, for example, corresponded to a wave profile that had a steep rising front and a long decaying profile after the peak. Blue light, in contrast, corresponded to a wave profile that began with a shallow, slowly rising profile followed by a rapid decline after the peak. Hooke's wave and color theory was never popular, and it is indeed physically incorrect, but it set the stage for more complete wave theories to follow – especially the one outlined by Christian Huygens in his 1703 text *Dioptrica*.

Moving Heaven and Earth

At the heart of Galileo's struggles with Church authorities was the heliocentric theory of Copernicus. The problem was not the theory itself, but the Church's insistence that Galileo should not publicly teach what he could not physically prove. Indeed, the Copernican model as originally published in 1543 had its problems – it was, in fact, fundamentally flawed with respect to its assumptions about the circular shape of planetary orbits. Using the planetary data collected by Tycho Brahe, however (who, in fact, didn't believe in the Copernican model), Johann Kepler was able to sort out the mathematical details and show that the orbits were actually elliptical, and that the speeds of the planets varied according to their position with respect to the Sun. Publishing his findings in 1609 and 1618 in two masterful works, *Astronomy Nova* and *Harmonius Mundi*, respectively, Kepler's three laws of planetary motion placed the Copernican model on a firm predictive footing.

As the seventeenth century unfolded so the acceptance of the Copernican model grew. There was no indisputable experimental evidence, however, to prove the actuality of Earth's spin and/or orbital motion about the Sun, and there was still the nagging problem that Kepler's laws had no known physical explanation. Confidence in the model was made absolute, however, with the publication of Isaac Newton's *Principia Mathematica* in 1687. This work, initially teased from a reluctant Newton by the ever diplomatic Edmund Halley, set science spinning in a new direction.

Newton's *Principia* was as revolutionary as Copernicus's *De Revolutionibus*, but rather than arriving to an unenthusiastic silence, his work was heralded with raucous acclaim. Newton had, with one text, changed the way in which science was to be done, and he established the rules and equations that described the motions of the planets. In one clean swoop, Newton explained why the Copernican model had to be right and why there should be three Keplerian laws of planetary motion – it all boiled down to gravity and the conservation of angular momentum.

Although Newton explained the theoretical side of things, the experimental verification of the Copernican model was still not established – an irritating leftover, one might say, concerning a now fully believed-in theory. Experimental evidence for Earth's motion was eventually found in 1725, however, after it had been observationally shown that the speed of light must be finite.

Dealing with Descartes

Much, if not all, of the *Principia* was written to satisfy Newton's dislike for the philosophy constructed by French savant Rene Descartes. Descartes was famous for developing a dualistic philosophy in which mind and matter were separated and distinguished according to the essence of thinking and spatial extension. All entities, and especially living ones, were thought of as machines, made up of inert and independent components, with an overarching order being imposed from within. This was the world of invisible ghosts within the machine (Fig. 2.5).

Descartes wrote about optics and light, but his reasoning is often contradictory and incoherent. Most of his ideas on light are contained within a *Treatise of Light*, which was written as part of a more extensive body of work published during the years 1629–1633. Descartes reasoned that light traveled instantaneously, as an impulse across adjacent particles, and yet allowed the speed of light to vary in different media. Indeed, Descartes never really makes it clear whether he thinks of light as being a form of motion or the movement of matter. All such philosophical wrangling aside, however, Descartes was able to reason his way to what is the correct mathematical form for the law of refraction – known today as Snell's law.

Snell's Law

Dutch mathematician and astronomer Willebrod Snell van Roijen derived his law for the refraction of light in 1621, although he never actually published it in his lifetime. Indeed, it was to be 82 years after its discovery that the law would first see print under Snell's name.

While Rene Descartes had published his derivation of the refraction law (identical to Snell's) in 1637, it was Snell who technically made the first innovations, according, that is, to Dutch scientist Christian Huygens, who publicized Snell's

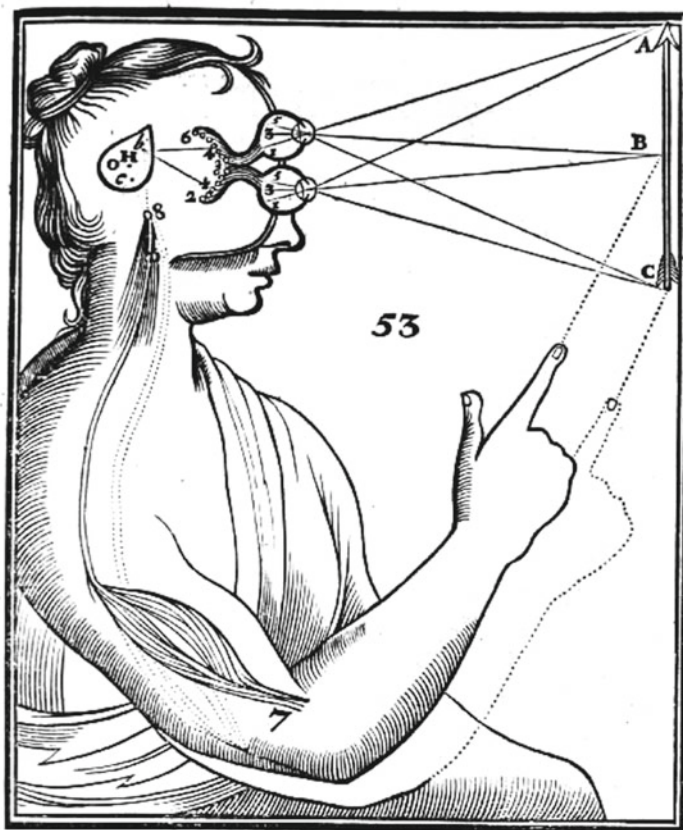


Fig. 2.5 An illustration depicting Descartes' proposed linkage between vision, which was interpreted by the brain and mind, and the mechanisms within the body responsible for reacting to external stimuli

work in his own 1703 masterpiece *Dioptrica*. Neither Snell nor Descartes, it turns out, can really claim priority for discovering the refraction law, since in all but the detailed mathematical language, Islamic scholar Ibn al-Haytham described the relationship in his *Kitab al-Manazir* (Book of Optics) written, while under house arrest in Egypt, during the decade 1011–1021.

The modern-day version of Snell's law is expressed in terms of the refractive indices of the two media through which the light is passing. The refractive index n is a measure of the speed of light within the given medium V_{medium} compared to that in a vacuum c (its fastest possible speed): specifically, $n = c/V_{\text{medium}}$ and this ratio is always going to be greater than or equal to one.¹

¹The refractive indices of various optical solids, liquids and gases are readily available in almost any compilation of physical constants.

We are now in a position to explain the term optical density introduced earlier with respect to the law of refraction (recall Fig. 2.3b). Indeed, the meaning is quite straightforward, in that the larger the optical density, the larger is the associated refractive index of the medium and the slower the light speed through it. The refractive index for air is usually taken as $n_{\text{air}} = 1$, although it is in fact variable according to the temperature and density of the air (as we shall see in Chap. 3). The refractive index of ice is 1.31; that of Perspex 1.495, while that of diamond is 2.417. The later result indicating, for example, that light travels nearly 2.5 times more rapidly through air than it does through a diamond.

With the definition of refractive index in place we can proceed to a statement of Snell's law. If the first medium has refractive index n_1 , and the second n_2 , then the relationship between the incident, i , and refracted ray, r , is given by the sine of these angles such that $n_1 \sin(i) = n_2 \sin(r)$. Snell's law is slightly more complicated than the law of reflection (recall Fig. 2.3a), but it is nonetheless a straightforward mathematical expression and simple enough to evaluate with an electronic calculator.

Fermat Finds the Quickest Way

Like Newton, French mathematician Pierre de Fermat was entirely unimpressed with Descartes theory of light, and he publicly said so. For Fermat this was a bad move. Descartes was incensed that his ideas should be so criticized, and he deliberately set out to destroy Fermat's reputation. Fermat, however, at least posthumously, receives the last hurrah since Fermat's principle is still taught in optics classes to this very day, whereas Descartes is rarely, if at all, mentioned.

Personal problems with Descartes aside, Fermat established the basis for what is now recognized as his principle in the mid-1650s. He built upon a classical foundation. Indeed, in the first century A.D. Hero of Alexander had noted that the law of reflection (that the angle of incidence is equal to the angle of reflection) could be explained upon the supposition that light traveled the shortest possible distance between the source and the observer. In the eleventh century Islamic scholar Alazen, took Hero's argument further and applied it to refraction – finding in consequence the forerunner to Snell's law. Fermat took this ancient idea one step further and argued in early 1662 that light travels not along the shortest geometrical path but along the path with the shortest travel time. This is a significant change in thinking and one that inherently adopts the idea that the speed of light is both finite and variable in different optical media. For Snell's law to hold true, for example, light must travel more slowly in glass or water than in air – in contradiction, in fact, to the claims of Descartes.

If Fermat's problems with Descartes were not enough to be getting on with, he sadly found no great support for his ideas on light either. The main problem was that his arguments were seen by others to imply knowledge-a-forethought to nature. How did a light ray *know* which was the path with the shortest time, it was asked, since, of course, light doesn't know and indeed it cannot reason.

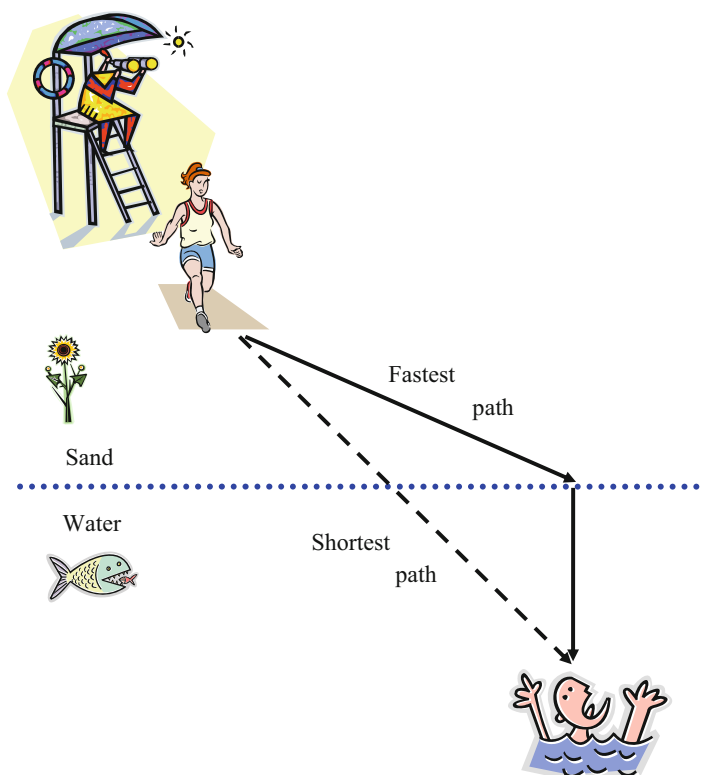


Fig. 2.6 Analog of Fermat’s principle. The lifeguard minimizes the time to reach the swimmer by minimizing the distance that she has to swim. The rationale for taking the crooked path is based upon the fact that she can run more rapidly on sand than swim through water, and thus the direct, straight-line path is not the one that has the least travel time

Philosophically, Fermat’s principle is not a problem if one views it simply as a mathematical means of calculating where a light ray will go – it is a means of deduction rather than an explanation of original cause. Eugene Hecht (Adelphi University, New York) has eloquently written before us, when it comes to Fermat’s idea, “let us embrace it but not passionately.” Not for the first time in physics, nor indeed for the last, are the means of calculating something about nature strange, spooky, and counterintuitive – we shall pick up on the problem of how light behaves, one more time, in Chap. 4.

There is a rather nice physical analog to Fermat’s principle that relates to the actions of a lifeguard and a swimmer in distress. Figure 2.6 illustrates the problem faced by the lifeguard, who is our proxy for the light ray – light goes where the lifeguard goes. Now, between the lifeguard station and the ocean is a stretch of sand-covered beach, and the lifeguard knows that she can run faster on sand than swim through water. So, in order to minimize the time in getting to the swimmer in distress she reasons that the shortest, straight-line distance between the swimmer

and the lifeguard station is not the optimal path to take. Rather she must take the path that minimizes the distance that has to be swum (because, remember, she runs faster than she swims). In this case the path taken by the lifeguard is the one that has a ‘kink’ at the sand-sea boundary, since this path minimizes the swimming distance. Fermat’s principle says that the path taken by a light ray can be calculated in the same manner as that carried out by the lifeguard. In modern optics Fermat’s principle is cast in terms of differential calculus and the determination of what are called stationary points in the optical path with respect to variations in path length.

Newton’s *Opticks*

Although the *Principia* came into print when Newton was a relatively young 44 years old, his equally provocative and revolutionary *Opticks* appeared in 1704, when he was a stately 61 years of age. Many decades in gestation, *Opticks* was a synthesis, and sometimes rambling account, of numerous ideas. Not least among its pages are Newton’s ideas on the properties of light and how it interacts with matter. The book also outlines his revolutionary ideas on how science itself should be conducted – building, as it were, on the twin pillars of mathematics and experimentation.

The story of Newton’s ideas on light begins in 1672, at which time he described his optical experiments to the members of the Royal Society in London. “In the beginning of the year 1666... I procured me a triangular glass prism, to try therewith the celebrated phaenomena of colour” – so wrote Newton in his preamble that led up to his *experimetum crucis*. Indeed, in this “critical experiment” he described how a glass prism could be used to show that white light must, in fact, be composed of many differently colored rays. Light, he argued, is a heterogeneous mixture of different colored components that the prism is able to separate out according to their specific ‘refrangibility’ (or, as we would now say, refractive index). To Newton the color of light was simply related to the size of its transmitting corpuscles. In 1675, in another letter to the Royal Society, Newton further explained that he supposed, “light is neither aether, nor its vibrating motion, but something of a different kind propagated from lucid bodies.” With this idea, that light can be envisioned as a stream of minuscule corpuscles, Newton was at odds with his fellow countryman Robert Hooke and Dutch scientist Christian Huygens, who advocated the more vibrant, wave theory of light. For these scientists, light was an oscillation that propagated through the “aether,” or ether, a diffuse medium that pervaded all space. While the propagation of his light corpuscles did not require an intervening aether, Newton nonetheless suggested that they did, somehow, interact with it.

Unlike Descartes and many of his contemporaries Newton advocated a finite speed for the propagation of light (indeed, Dutch astronomer Ole Romer, as we shall see momentarily, had shown it must be finite in 1676). In addition to light speed being finite, Newton conjectured that its speed would vary according to the optical properties of the medium through which it was passing. He reasoned, for example, that the speed of light would be slower in air than in glass or water since in the latter,

because of their higher densities, there was a reduced quantity of aether. This was a crucial distinction between Newton's theory and its competitors, which suggested that light should move more slowly through more dense optical media. Indeed, as we shall see shortly, the speed issue was the Achilles heel of Newton's theory.

The Inconstant Moons of Jupiter

Galileo first noticed the four moons that orbit Jupiter on the night of January 7, 1610. It was an incredible find, and Galileo knew he could profit from it; scientific fame was assured, and with the right dedication to a rich patron a better paying job was his for the taking. Writing of his telescopic discoveries in *Siderius Nuncius* Galileo suggested that the newly discovered worlds be called the Medician stars.² The Medician family was impressed by the gesture, and Galileo began to climb the slippery and more often than not dangerous slopes of the Italian Renaissance social ladder.

Now officially named Io, Europa, Callisto, and Ganymede the Galilean moons are classed among the most incredible and fascinating objects within our Solar System. Io is a volcanic world, continually turning itself inside out; Europe has a global ocean underneath an icy crust; Callisto also supports a hidden sea. Ganymede is a behemoth of a moon, and in fact outsizes the planet Mercury by several hundred kilometers. Io, the innermost of the Galilean moons, takes some 42.456 h to complete one orbit about Jupiter. Further out and therefore moving more slowly Callisto takes 16.689 days to traverse its orbital path. The moons tap out a stately round and furnish the heavens with the hands of a pointillist clock – provided, that is, we can work out how to read the message on display.

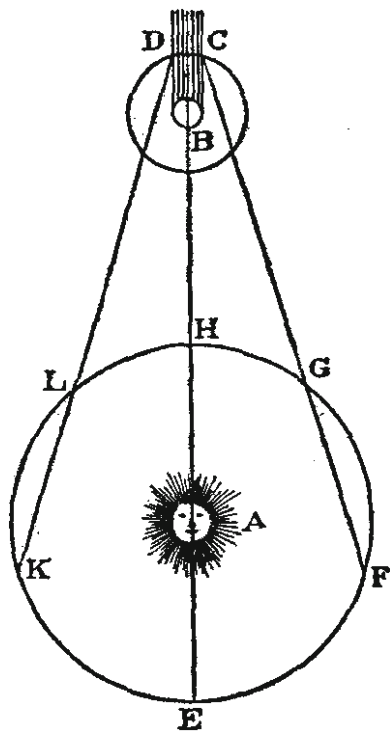
Galileo correctly reasoned that if tables of the exact transit and eclipse times of the Medician stars could be constructed from one fixed, specific location then navigators and observers at different locations could use these tables and their local time observations of the moons to determine their longitude.

Systematic observations of the Jovian moons began soon after Galileo publicized his discovery, but it was Italian mathematician Giovanni Domenico Cassini who first noted, in 1675, something puzzling about their motion – the apparent transit and eclipse times of the moons varied according to the distance that Jupiter was away from Earth. From this observation Cassini correctly reasoned, at first, that this was a direct consequence of the finite speed of light. He later changed his mind, however, and threw in his lot with the philosophy outlined by Rene Descartes and reasoned that light must travel instantaneously.

Taking a second look at the problem, however, Danish-born astronomer Ole Romer, working at the Observatory in Paris, made precise timing measurements and

²Galileo's suggestion was never popular. The satellites (not stars) were and still are invariably known as the Galilean moons of Jupiter.

Fig. 2.7 Romer's diagram for explaining the maximum occultation time advance and delay (compared to an imagined observer at A) for a Jovian moon. According to the location of Earth being at either E or H the eclipse times will be early and late, respectively. The speed of light can be deduced by determining how long it takes light to cross the distance HAE



concluded that the eclipse times of Jupiter's moons truly did indicate that light traveled at a fast but finite speed. Remarkably, upon announcing his results at the November 21, 1676, meeting of the *Academie des Sciences* his most vociferous opponent was Cassini.

What was it that Cassini and Romer observed? The geometry of the situation is illustrated in Fig. 2.7. The circle LHGFKEK indicates Earth's orbit around the Sun (located at A), and B indicates the position of Jupiter. The circle around Jupiter represents the orbit of one of the Galilean moons. The points C and D illustrate the locations at which the particular Jovian moon will begin and end an occultation, and these are the observable events that can be timed from Earth. What Cassini and Romer found was that compared to an imagined observer located at the Sun (A), who would see occultations and transits at equal intervals of time, the timing of the same events observed from Earth would be either early (when Earth was at position H, for example), or late (if Earth was at position K or F). The time difference, being early or late, correlates with Earth being closer or further away from Jupiter. The maximum time delay between the early and late occultation times will be that recorded between points H and E on Earth's orbit. And, this maximum time delay will correspond to the time it takes light to cross Earth's orbital diameter. In modern terms, Earth's orbital diameter is known to be about 300 million kilometers across, and it will take light some 1,000 s (or about 16.5 min) to cross this distance.

While the observational results presented by Romer proved that light traveled at a finite speed, he deferred from actually calculating its speed since at that time the physical measure of Earth's orbit was not well known. Finally, however, after more than 2,000 years of speculation, the speed of light had been brought down to size, no longer infinite, but certainly extremely fast.

The Aberration of James Bradley

The reverend James Bradley (1693–1762) was a serious if not intimidating man. In many ways, however, he can be considered to be the founder of precision astronomy, and he certainly had a penchant for finding small angular corrections. As an astronomer he spent many nights during 1725 lying down – not sleeping, of course, but carefully peering through the lens of a long zenith telescope. The work was painstaking, exacting, and ultimately revolutionary. He was looking for a small semi-annual shift in the position of one specific star – the star gamma Draconis; also known under the Arabic name of Eltanin (meaning *Serpent*). The star had been chosen quite deliberately since it passes almost directly overhead for an observer stationed in London. By passing through the zenith the star's position could be measured with exquisite accuracy, and what Bradley was hoping to measure was a parallax effect due to Earth's motion in its orbit around the Sun over a 6-month time interval. In this manner Bradley was working at the cutting edge of eighteenth-century astronomy, and he was also hunting for one of astronomy's holiest of grails. If he could measure the parallax of Eltanin then he could calculate how far away it was, determining thereby, for the very first time, the scale of interstellar space. Bradley failed in his parallax quest, but he succeeded in finding something else that in retrospect was just as important. He proved that Earth must be in motion about the Sun.

Over the course of a year's observations Bradley tracked and accurately measured the position of Eltanin as it passed overhead. He found a 6-monthly shift of 0.01° . There was a problem, however, since the observed direction of the shift was not that which would be expected for parallax. It was an unexpected aberration – indeed, it was stellar aberration. What Bradley had discovered was that Earth must be in motion around the Sun (as Copernicus had said it had to be, and as everyone then firmly believed anyway), but now there was experimental proof. Galileo would have been elated – Copernicus had finally been vindicated by direct experimental measure.

The effect that Bradley had observed can be seen during any rush hour scramble on a windless, rainy day. As the pedestrians scurry to work or run for their buses, you will notice that their umbrellas are not held vertically but are distinctly angled forwards in the direction that they are moving. The umbrellas are tilted forward, not because of the wind, but so that the vertically falling rain doesn't soak their lower body and legs. People angle their umbrellas unconsciously to compensate for their motion through the rain; astronomers must do essentially the same thing and tilt

their telescopes in the direction of Earth's motion, as it moves around the Sun, in order to accurately track the position of a star.

Bradley's observations, in addition to revealing Earth's orbital motion, also allowed for a refined measure of the speed of light to be made. In a 1729 letter to the Royal Society Bradley announced his findings – the speed of light was no less than 185,000 miles per second (or about 298,000 km per second). It was an incredible result, and one that was not to be bettered by laboratory experimentation for another 120 years. Bradley's fame was in its ascendancy, and in 1742 he was awarded a plumb job, becoming Britain's third Astronomer Royal.

As we move into the early nineteenth century, it had become clear to most scientists that light must travel at an astounding speed, and that it behaved as if it were a wavelike phenomenon. But, what is a wave and how can waves be quantified? This we shall discuss next, and thereafter describe the definitive experiments that seemingly put paid to Newton's tiny corpuscular rays once and for all.

Waves and Wave Interactions

Waves and wavelike phenomena surround us. The world literally undulates its way toward the future; the Sun rises and the Sun sets (a day-night wave), earthquakes cause the ground to tremble (surface seismic waves), and traffic on the way to work each morning stops and starts (a frustration wave?).

Waves can have many different sizes, shapes and forms, and can mix and meld. For all this variation, however, there is a remarkable theory relating to waves that was first derived by French mathematician Jean Baptiste Joseph Fourier during the second decade of the nineteenth century. Fourier showed that all periodic waves, no matter what their profile looks like, could be constructed by adding together a series of simple sine waves, each of a different amplitude and frequency (these terms are defined below). The basic unit of all waves, therefore, is the sine wave.

A sine wave is the sort of wave profile that we have all seen at one time or another. A wave rippling along a rope held between, say, two white-coat clad experimenters would be an example of a sine-like wave. With this image of a rippling rope in mind, we also learn something else about waves – a wave is a disturbance that moves through a medium (the rope in our image). That is, the wave describes the time and space deformation of the medium away from an equilibrium (or rest) position. The rope, in our specific case, periodically moves up and then down about a horizontal line (the rest position) at each and every point along its length. The wave disturbance moves along the rope, but the rope itself does not physically move between the two experimenters – all the motion in the rope is in the vertical up and down sense.

A perfect sine wave has four essential characteristics: its wavelength, its period, its amplitude and its propagation speed. The wavelength is the spatial distance, λ , over which the wave profile starts to repeat itself (Fig. 2.8), the period is the time taken to repeat the profile and the amplitude is a measure of the displacement of the wave away from the equilibrium (or zero) point. The time characteristic of a wave

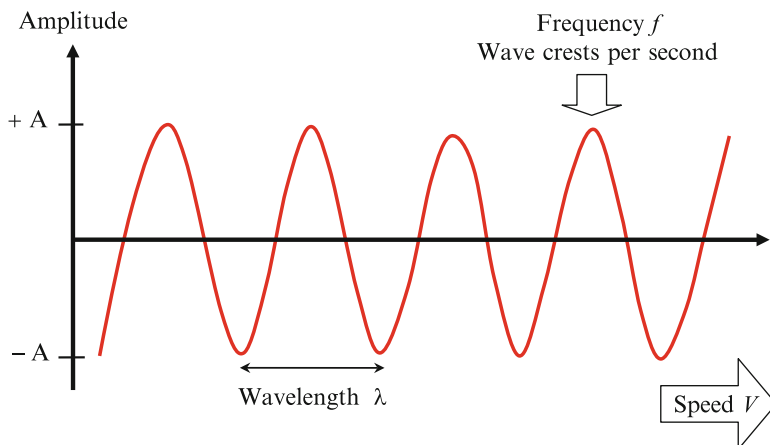


Fig. 2.8 Characteristics of a wave. The frequency is related to the number of wave-crests that pass a specific reference point (*vertical arrow*) per second. The speed of the wave is the product of the wavelength times its frequency

is usually expressed in terms of its frequency, f , which is just the inverse of its period. The speed V with which a wave propagates through a specific medium is related to the wavelength and frequency according to the formula: $V = \lambda f$.

One of the most important characteristics of waves is that they can be added together, or superimposed, to form another wave. The resultant wave, however, can look very different from the original pair, according, that is, to their so-called phase difference. If two waves are exactly in phase then the amplitudes of each will add at every instance and the resultant wave will have twice the amplitude at every point in its profile. This is called constructive interference.

If, on the other hand, the two waves are combined such that one wave is shifted one half wavelength on from the other (a phase difference of 180°), then the peaks and troughs of each wave will cancel each other out, and in fact no wave will result – the summation will be a straight line. This is destructive interference. As we shall see in Chap. 5, destructive interference is one method by which aircraft can be made invisible to radar waves. The two cases of interference as just described are illustrated in Fig. 2.9. In general the phase difference can be any fraction of a wavelength – or fraction of a period – and the resultant wave summation, that is, interference, can be quite complex, as we shall see later on.

The Triumph of Waves

With the finite speed of light established, the next problem that the natural philosophers had to sort out was the question, “What is light?” Is it a wavelike phenomenon or is it a stream of small corpuscles? Newton had made his thoughts clear in 1704: light

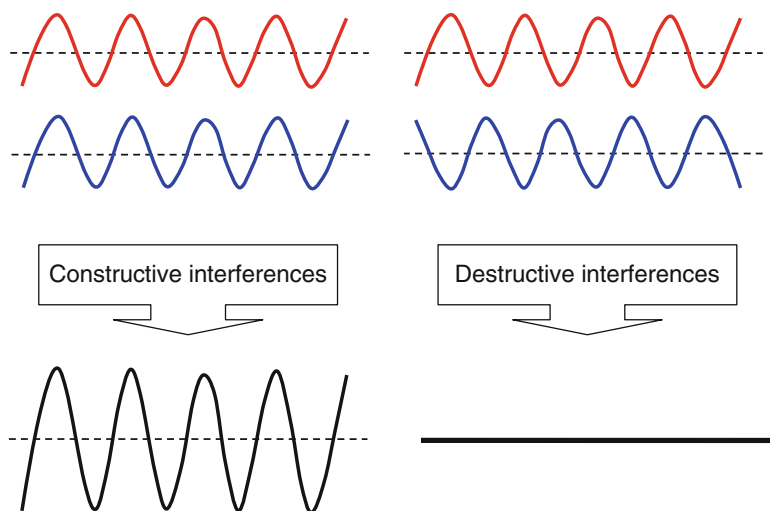


Fig. 2.9 (a) The constructive interference of two waves. (b) The destructive interference of two waves

was composed of a stream of different-sized particles. And, while Robert Hooke had advocated an alternative wavelike theory for light, it was Christian Huygens in Holland who really established the detailed working theory.

Writing in his 1678 work *Traite de la Lumiere*, Huygens articulated his principles of wave optics. Light could be thought of as a traveling wave front that moved through space. The idea here is that we can imagine the light to stream away from an illuminated source with successive wave fronts spreading out one after the other as ever-growing spherical boundaries – just like the ripples moving over the surface of a mill pond after a stone has been dropped through its surface at some point. Importantly, Huygen’s also introduced the idea that every point on an advancing wave front could be thought of as the potential center for a new disturbance capable of generating its own specific wave front. An everyday analog of this principle is demonstrated by sound waves.

Imagine that we have two rooms that are connected by a single small doorway. In one room is an observer – or listener – and in the other room is a sound source, say a radio. Everyday experience, or if you are not convinced experimentation, now tells us that no matter where the radio is located within the one room – in a corner, in the middle, on the floor, or on the ceiling – the listener in the adjacent room will always hear the sound as if the radio is located at the door itself. (This analog, as all analogs invariably are, is limited and assumes there is no sound conduction). The same behavior as displayed by our analog sound waves should, according to the Huygens principle, also be present in light, and in 1803 the British polymath Thomas Young demonstrated that this was indeed the case. Young’s slits experiment is a classic, and it is still a standard right of passage to this very day for any student taking a college or university physics program. It is a beautifully simple experiment, and yet it delivers a result of profound importance.

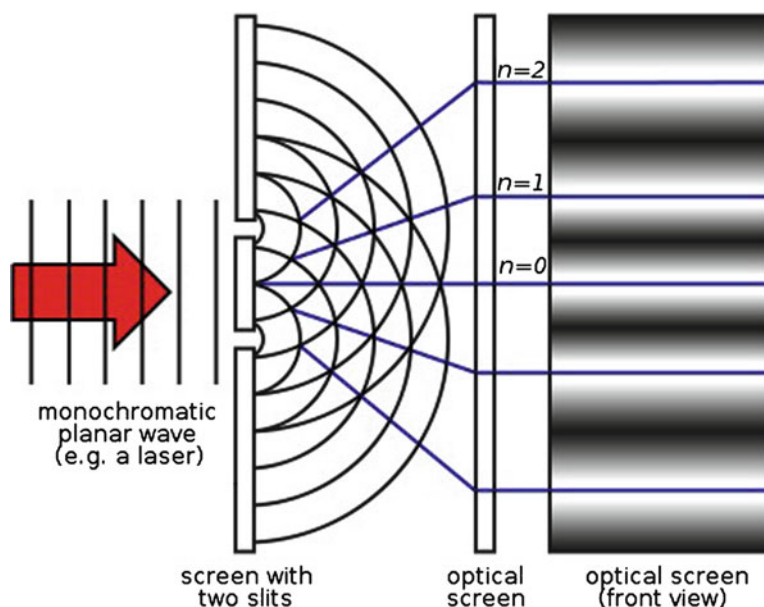


Fig. 2.10 Young's slits experiment for the demonstration of the wavelike behavior of light (Image courtesy of theskepticsguide.org)

The essential layout of Young's experiment is shown in Fig. 2.10. The key idea behind the experiment is to set up the conditions for constructive and destructive interference (recall earlier in this chapter). With respect to Fig. 2.10, we can imagine a parallel wave front approaching a panel in which two slits have been cut. Each of the slits in the panel will accordingly act as the source of a semi-spherical wave front, but now the expanding wave fronts will begin to cross over each other, setting up the conditions for interference. Where two wave fronts interact in a peak-to-peak manner the amplitude will be enhanced (constructive interference), and where two wave fronts interact in a peak-to-trough manner the amplitudes will cancel to zero (destructive interference). The ultimate result of the interference between the expanding wave fronts will be a series of light and dark bands being formed on a screen (to the far right in Fig. 2.10). Upon actually seeing such interference bands on his experiment's screen, light, Young concluded, must be a wavelike phenomenon.

Faster or Slower? A Crucial Test

If, with the publication of Newton's *Principia* in 1687, French scientific pride had been deflated through the sound trouncing it gave to Descartes' plenum-packed universe of mystical whorls and eddies, so the savants of Paris were able to regain some revenge by putting Newton's optical theory to the test. The predictions between

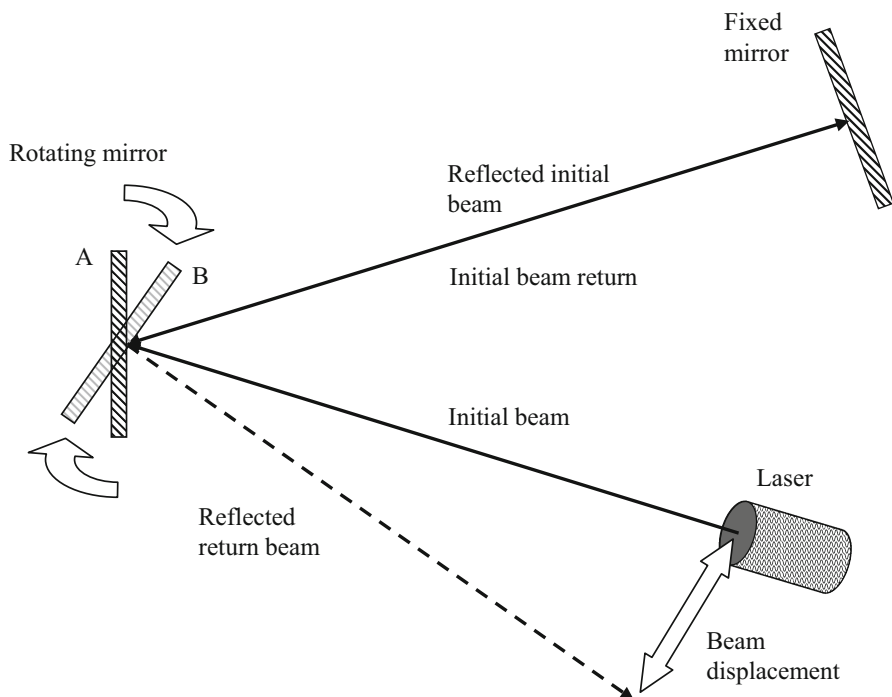


Fig. 2.11 Schematic outline for the rotating mirror experiment used by Leon Foucault to determine the speed of light in air and in water

the rival theories were, for once, quite clear. The wave theory required that the speed of light must decrease in a denser optical media (such as glass or water), whereas the corpuscular theory of Newton (and Descartes for that matter) required that it travel faster. It would be a case of winner takes all.

If an exact experimental measure of the speed of light through water, say, could be found, or even just a difference in its velocity recorded then the critical question could be answered. French experimental physicist Leon Foucault was the person up to the challenge, and at the heart of his experimental test was a whirling mirror (Fig. 2.11).

The beauty of Foucault's test was that it measured a relative response rather than an absolute one. The idea was not to physically measure the speed of light in water, but to see what affect a column of water had upon the position of an imaged light ray. In a modern-day version of Foucault's experiment (as in Fig. 2.11) a laser sends out a thin pencil beam of light that is directed in order to reflect off a rapidly spinning mirror. Upon being reflected from the surface of the rotating mirror the light beam is directed towards a distant stationary mirror. After reaching the fixed mirror the light beam is halfway around its path and is reflected back towards the first, still rapidly rotating mirror. When the light beam, however, encounters the rotating mirror again, it will have moved slightly around its spin axis and accordingly the light beam, after a final reflection, ends its path being directed towards a screen displaced

away from the laser. Technically, if the distances between all the mirrors and the lamp and screen are measured, and the mirror's spin rate is known, then the speed of light can be determined, but this is not the main task of this particular experiment. The point, remember, is to see if the speed of light increases or decreases during its passage through a water column.

The experiment is now half complete. The position of the returned light beam on the screen is carefully recorded – in practice this is done with a microscope and a finely divided scale. The trick now is to place a column of water (in a sealed glass tube) in the beam path somewhere between the rotating and the fixed mirror. The question then becomes, what happens to the returned beam displacement? If light travels more rapidly through water, as Newton's corpuscular theory predicted, then the returned beam will hit the screen closer to the laser than determined under the open-air configuration. This is because the light beam will hit the rotating mirror on its return path sooner than before, and the mirror will have accordingly spun through a smaller angle and the final reflection will also be at a smaller angle than in the open-air case.

The exact opposite effect will occur if light travels more slowly through the water column. In this case, the returned light ray will take longer to reach the rotating mirror, which will have spun further around its axis, and accordingly the final reflection will intercept the screen at a larger angle of displacement than that observed in the open-air configuration. It is a brilliantly simple experiment.

Foucault fussed over the experiment but in 1850 was ready to release his findings. The water column caused the light beam to return later and the displacement was away from the laser and the in-air only reference position on the screen. Light travels more slowly through water, and Newton was wrong – the corpuscular theory had been trounced, and the wave theory advocated by Huygens and Hooke vindicated. Light must be a wave.

But with one theory eliminated another problem arose. What was it that was 'waving,' and through what kind of medium did it propagate? These were not, however, Foucault's problems to solve, although in later life he came back to refine and finesse his experiment and was able to measure ever more precise values for the speed of light.

His next experimental triumph, in 1851, was another scientific revolution. The revolution in this case was the first experimental demonstration of the spin of Earth with a giant pendulum suspended from the inside of the dome of the Parthenon in Paris. This verified experimentally the rotation theory advocated by Copernicus in 1543.

At this stage, Foucault drops from our story and the next major player we encounter is James Clerk Maxwell. As we shall see in Chap. 4, Maxwell was one of the scientific giants of nineteenth-century science – his mark is indelibly stamped on the foundations of many branches of physics, ranging from thermodynamics to optics and classical mechanics, and to the development of electromagnetic radiation theory. It was Maxwell who, in the mid-1860s, forged a link between magnetism and electricity – fields that were previously thought to be distinct. And, by forging this linkage he was able to show that light is a double wave (Fig. 2.12). The phenomenon that we call light is, in fact, a waveform with both a varying magnetic and

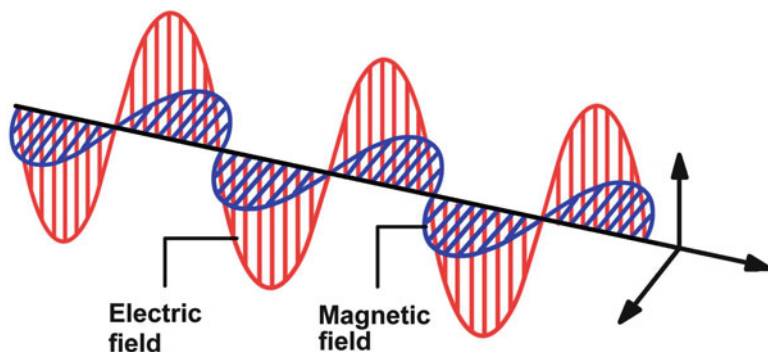


Fig. 2.12 The double waveform of light. The magnetic and electric fields oscillate at right angles to each other

a varying electric component, with the two components oscillating at right angles to each other. Maxwell's analysis was new and stunning, but it took some time, nearly 30 years in fact, before it was fully realized just how stunning and useful a result it really was.

Polarization

The polarization state of a light beam is, by convention, described according to the orientation of the plane in which the electric field vibrates (Fig. 2.12). The measurement is in the plane perpendicular to the direction of travel, and if the wave crests are all orientated in the same plane then the light is said to be linearly polarized. The plane of the electric field can also continuously rotate as it moves along, and in this case the light is circularly polarized. Not all light is polarized. Starlight, for example, is produced by a large number of randomly moving atoms, and there is no correlation between the orientations of the resultant electric field plains – the light is unpolarized.

The French mathematician Etienne Louis Malus first described the polarization of light in 1809, and it turns out that light is generally polarized in the process of reflection. The light from the Moon, for example, is polarized since the illuminated fraction that we see from Earth is, of course, reflected sunlight. Light seen in reflection from the surface of a lake or a wet road is also polarized. Indeed, by adding a polarizing component to spectacle lenses irritating glare can be reduced and image quality greatly improved. How does this work?

The way in which polarizing filters, or polarizers, work is to force the electric field orientation into one specific plane – that is, after passing through the filter the transmitted light is linearly polarized. Essentially the material out of which the polarizer is made preferentially absorbs in all but one specific direction (the polarization axis). Longer wavelength radiation, such as microwaves (to be described in

Chap. 4), can, for example, be linearly polarized by making them pass through a closely spaced grid of parallel conducting wires. Indeed, an analogy that is often used to describe polarization is that of a person trying to send a wave along a rope that at some point passes through the gap in a picket fence. The person trying to “flick” the wave along the rope must send the wave in a direction that is parallel to the boards of the fence, else the rope will physically hit one or both of the wood panels and the wave will be destroyed.

The linear polarization properties of sunglasses can be seen if a light source, say a light bulb, is viewed through two pairs of glasses. By rotating one pair of sunglasses in front of the other the brightness on the image will vary from full brightness (when the polarizers are parallel) to no image at all, that is, no transmission of light, when the polarizers are at 90° to each other. Indeed, the variation of brightness is described according to the angle between the two sets of polarizers. (It was, in fact, this very law that Malus described in 1809).

Polarization measurements can be used to reveal otherwise invisible characteristics of many objects. Indeed, engineers often construct plastic models of new machine parts or architectural assemblies so that they can be subject to controlled stresses and loads – the point being that the applied forces will change the polarization properties of the plastic and thereby reveal regions where failures might occur.

Astronomers also measure the polarization of starlight to map out the magnetic fields that wind through galaxies. This latter situation is a remarkable way of making the inherently invisible visible. Starlight itself is unpolarized, but it turns out that the interstellar medium is pervaded by miniscule but slightly elongated dust grains. These dust grains can become orientated in a direction perpendicular to galactic magnetic field lines (we shall describe the meaning of field lines in the Chap. 4). Since the grains are elongated and aligned with the magnetic field, however, they are able to scatter and partially polarize starlight. The polarization angle of the starlight can then be measured and the magnetic field line structure of the galaxy revealed (Fig. 2.13).

What a remarkable chain of circumstances this truly is – an invisible magnetic field revealed by invisible dust grains interacting with distant starlight to determine the plane of oscillation in an electric field – science fact is indeed stranger than science fiction.

Sight and Sound

The early members of the Royal Society of London, founded in 1660, were, by modern standards, a bloodthirsty lot. They liked nothing better, for so it seems, than performing experiments involving animals – cutting them up, giving them blood transfusions, and killing them off in vacuum chambers. Well, of course, times were different then, and there was sometimes a serious point being made. The experiments with the vacuum chamber are particularly poignant to our story.

The vacuum pump was the “big science” instrument of the seventeenth century. Used by Robert Boyle, for example, to investigate air (or to be more precise regions

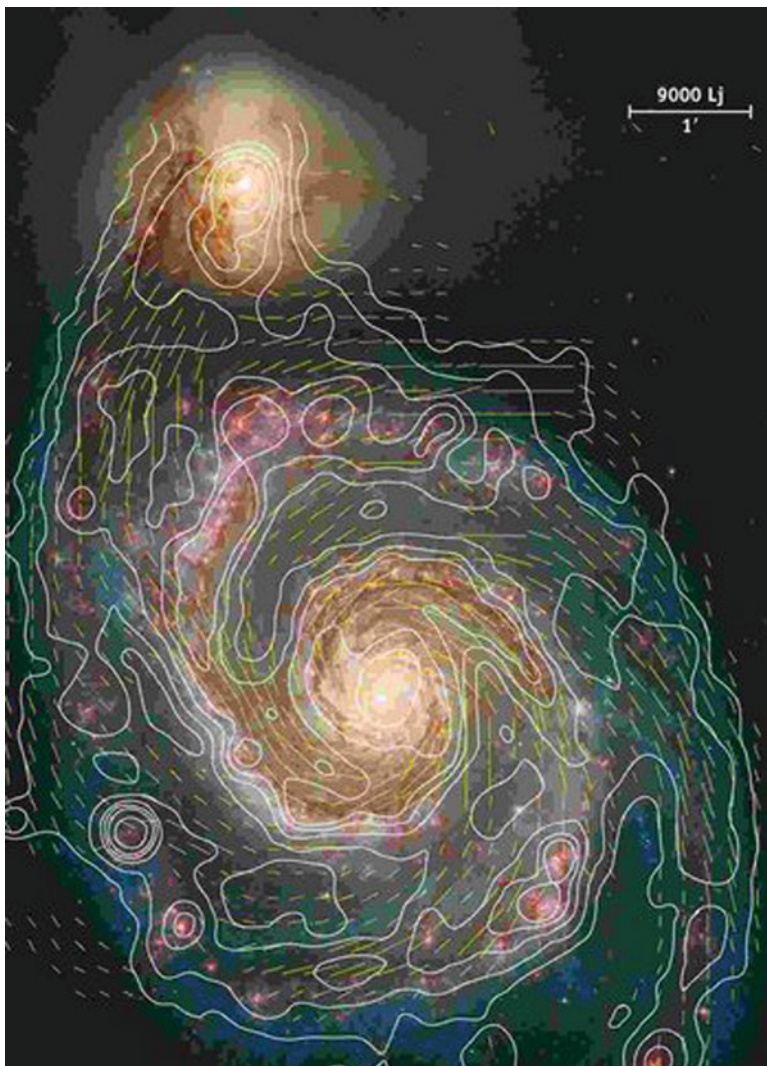


Fig. 2.13 Composite image of the aptly named Whirlpool Galaxy. The optical image captured by the Hubble Space Telescope has contours of constant radio emission intensity superimposed upon it. Also shown are a series of short lines indicating the direction of linear polarization and hence, by default, the direction and shape of the galaxy's magnetic field (Image courtesy of NASA)

were air wasn't) the early vacuum pumps were the technical marvel of the age (Fig. 2.14). The investigations performed by Robert Hooke with the air pump included experiments to see what would happen to burning candles, as well as small birds and animals, when the air was removed from the glass jar in which they had been placed. Other researchers, risking life and limb, experimented upon the firing of gunpowder within vacuum jars.

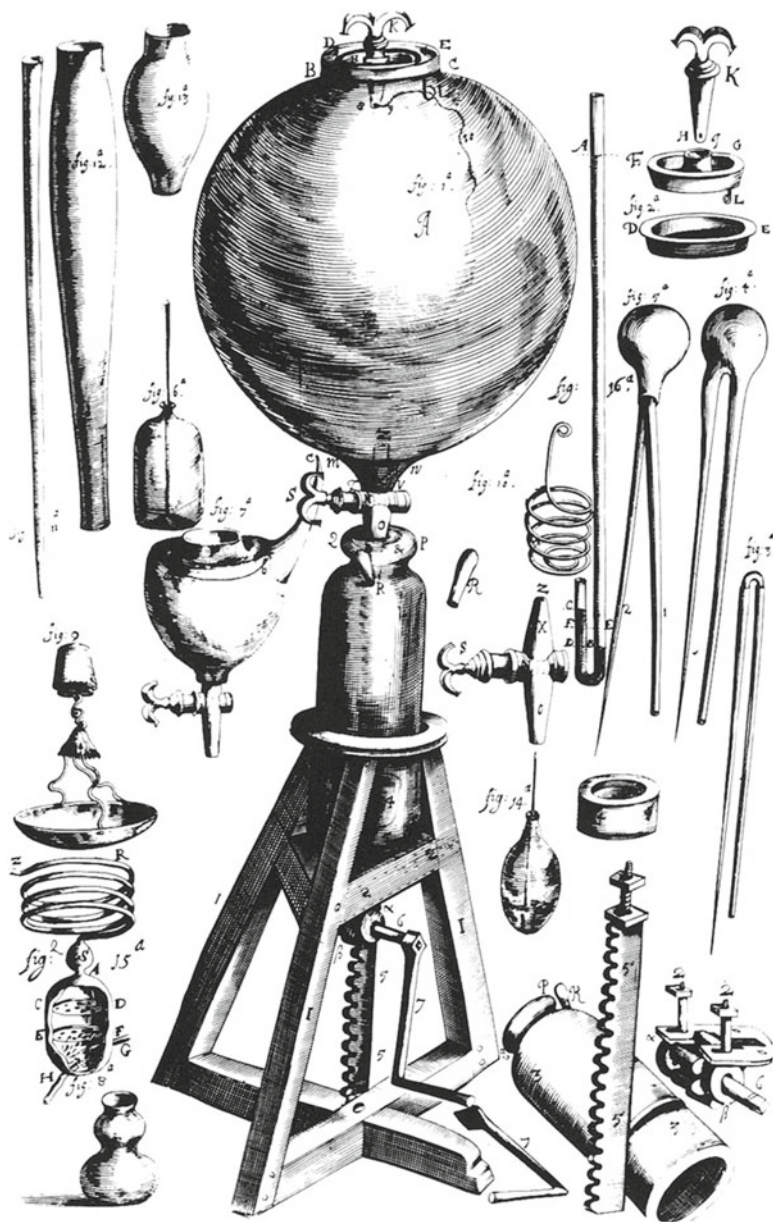


Fig. 2.14 Robert Boyle's vacuum pump, as perfected by Robert Hooke (Image courtesy of the Royal Society of London)

It was soon clear that both life and combustion needed something contained within air for their sustenance – and that something was eventually identified as oxygen. Later observations also revealed that sound waves could not propagate through a vacuum, demonstrating that it must be a pressure wave.

The propagation of sound experiment is now a classical high school science demonstration: an electric bell is placed within an inverted bell jar, and a vacuum pump is used to slowly extract the air. As the air is removed so the sound of the bell becomes weaker and weaker, eventually falling into an eerie silence.

Although this experiment tells us many things about sound waves, it is easy to miss another, perhaps more fundamental, observation. While we no longer hear the bell once the air has been removed from the vacuum jar, we can still see the bell and its parts vigorously vibrating. Light must clearly be a very different wave to sound, since light can propagate across a region in which the air has been removed. How does light achieve this remarkable trick? The answer to this question will not be revealed until Chap. 4.

In this chapter we have reviewed, albeit very briefly, the history of light. We have met some of the key players involved in explaining its properties, and we have looked at a few of the ways in which light has been experimented upon. Indeed, by learning how to weave light between lenses, mirrors and polarizers scientists have discovered two new worlds – the world of the very large, now surveyed with gigantic telescopes, and the world of the very small, scanned with ever more powerful microscopes. These optical devices extend our intellectual reach to dimensions and exotic domains otherwise invisible to the unaided eye. We are now ready, and hopefully enlivened, to take an initial look at a few of the ways in which light can be manipulated, diverted, scattered and tricked – and indeed engineered to make solid objects disappear from view.

Chapter 3

Mirror Tricks, the Art of Cloaking and Seeing the Impossible

Smoke and Mirrors: Without the Smoke

Look into any flat mirror. What do you see? Well, presumably, your face. Now, put your hand in front of your face. What do you see? This time, of course, you see your hand. More interestingly, we might now ask, is it possible to arrange a series of mirrors so that, once again, you'll see your face – making your hand, as it were, invisible. The answer to this question is, in fact, yes, and the means of performing this slight-of-hand is to construct a mirror ramp made of two periscopes (Fig. 3.1).

The periscope is probably best known through its military use in submarines and surveillance exercises where the observer wishes to remain unseen. The device is simplicity itself and is composed of two mirrors, one placed directly above the other, and both set at an angle of 45° .¹ In this fashion the light ray path is made to shift through a vertical distance equal to the separation between the two mirrors.

Four flat mirrors will be required to complete a mirror ramp cloaking device. As shown in Fig. 3.1 the idea is to construct two mirror-corners. One mirror-corner set will have the mirrored surfaces on the upper side, while the other has the mirror surfaces on the downward-facing side. A horizontal light ray will therefore undergo four reflections, each through 90° , such that after the fourth reflection it is moving in a horizontal direction again. Given this configuration anything placed under the lower mirror-corner will be invisible to the observer. The system is far from ideal, of course, and will only work if the observer looks straight at the mirror; a sideways glance will produce a distorted image that will immediately reveal the presence of extra mirrors. In addition to this, the actual mirror support struts will also be very obvious to any observer.

A slightly more refined two-dimensional cloaking device could be constructed with just three mirrors: in this case a small flat mirror is placed above two curved

¹ The Science Museum in London, England, provides the design information for a rotating, 360° periscope. See: http://www.sciencemuseum.org.uk/educators/classroom_and_homework_resources/resources/360_degree_periscope.aspx. (Page accessed October 2011).

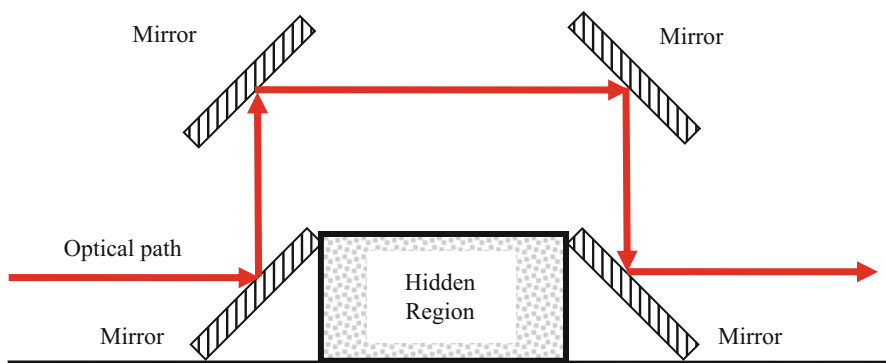


Fig. 3.1 A simple two-dimensional invisibility device. The object under the mirror ramp will be hidden from the direct view of the observer provided, that is, that he or she looks directly at the mirrors

mirrors arranged back to back. The two curved mirrors make the ramp under which an object can be hidden from direct view. In this configuration it is important that the curved mirror surfaces are parabolic in profile. The curved, specifically parabolic, profile is used since it will direct horizontal rays of light to the same focal point, where the small flat mirror is located in order to reflect the rays back to the second curved mirror in the ramp.

After this third reflection the rays will be moving horizontally again and in the same direction as the original path. Once again the device will only work if the observer is looking right at the mirrors, and any elevated view will reveal the presence of the cloak and the object being hidden. So far, therefore, we have only a selective cloaking device; the far harder, but not impossible, task is to construct an invisibility shield that will work from any viewing angle – as we shall see.

The Infinity Box

By everyday experience we are all familiar with the observation that a single plain mirror produces a single image of any object placed in front of it. By using two mirrors facing each other, however, one in front of your head, say, and the other behind it, the number of images can, in principle, jump to many tens (if not technically to infinity). Here we find ourselves with an analog of the regression to virtual infinities in both directions, placing our eye at the nightmare abyss that so frightened French mathematician Blaise Pascal. Indeed, we live at the junction between the infinitely small and the infinitely large.

Infinity mirror boxes are very simple to construct. The aim is to set two mirrors so that their reflective surfaces face one another. Any size of mirror can be used, but a small hole just a few millimeters in diameter must be carefully drilled through the center of one of them – this will be the peephole for the eye. The peephole is perhaps more conveniently formed with a pin pushed into a card surface that has had one

side coated with mirror tape – glass mirrors can be drilled, but this is a delicate process. With the peephole formed, arrange the mirrors so that they face each other and are separated by a distance that is about the same as their diameter, or according to shape their longest dimension. Now, suspend a small object, or even place a low-powered LED light, between the two mirrors and view the scene through the peephole. The image of the object will be repeated again and again to the eye – a veritable phantasmagoria of images will appear before you, not by magic, of course, but through multiple reflections of reflections of reflections and so on ad infinitum.

The infinity mirror box not only takes us to the teetering edge of a double infinity, it is in principle one way in which an object can be concealed while situated in plain view. In this case it might be said that the object is made invisible by multiple confusion. While it is true that the object is actually seen, it is not necessarily clear to the observer which is the real object and which are its multiply reflected doppelgangers.

Pepper's Ghost

Rather than working to disappear into the background the theatrical effect known as Pepper's Ghost² operates in the opposite sense, making an image appear on a carefully located transparent screen. The screen, placed between the actors on stage and the audience, allows for the projection of ghost-like figures into a play; it is very much a trick of the light. The arrangement of stage, actors, audience and projection system used to produce a Pepper's Ghost are shown in Fig. 3.2. The key optical component is the 45° angled sheet of glass situated at the edge of the stage. The audience can see the actors through the glass plate at the same time that an image of an under-stage actor (the ghost) is projected onto it, giving the appearance of an on-stage apparition. The glass sheet essentially works as a partial mirror, allowing light to pass directly through it (from the stage) and light to be reflected off its front surface (from the sub-stage).

Pepper's Ghost is a theatrical device by which two independent scenes can be optically joined, reminding us of the fact that what we think we see, as an audience member in this case, is not always a true representation of what is going on before us. A modern-day example of a Pepper's Ghost being put to work in a utilitarian fashion is that of the Heads-Up Display found in aircraft cockpits and more recently in some automobiles. In the former case, such devices project an image of important flight information onto the cockpit window, and this enables the pilot to see both the aircraft status data as well as the surrounding airspace. Pepper's Ghost is also very much alive and well, and indeed of central importance to a number of augmented reality devices and invisibility cloaks that have been recently developed in Japan.

² While developed for the stage by British physicist John Henry Pepper in the early 1860s, the idea was first described by engineer and part-time inventor Henry Dircks, and the effect was originally known as the Dircksian Phantasmagoria.

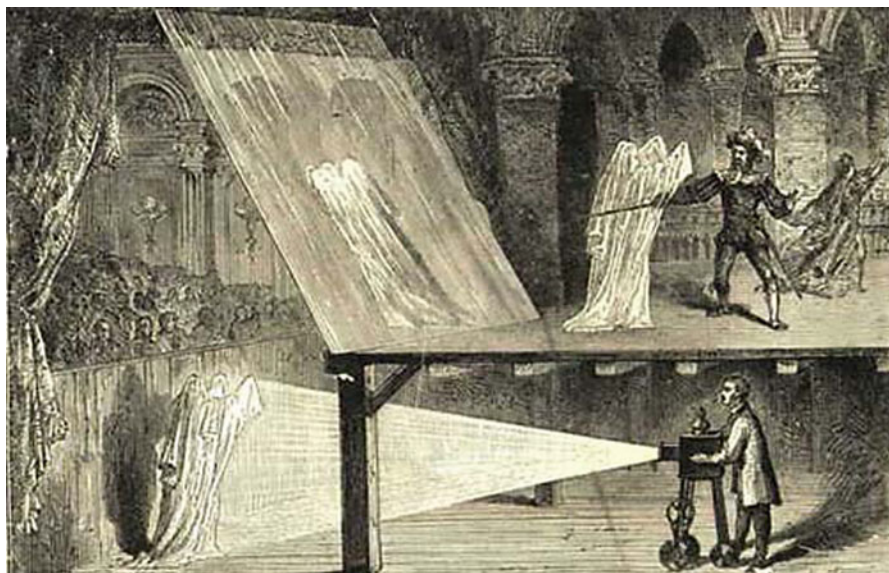


Fig. 3.2 The glass panel arrangement required to produce Pepper's Ghost. The ghost actor is strongly illuminated out of sight of the audience (*below stage*), and it is their reflection, channeled by the glass plate, which gives the appearance of a ghostly figure on stage

The Tachi Invisibility Cloak

One of the critical problems associated with mirror ramp invisibility devices is that mirrors are difficult to mount and arrange correctly (as you may have found in the infinity box experiment). Any small deviation from the true angles of alignment will produce telltale distortions readily visible to the observer's eye. Not only are these devices not easily transportable, but mirrored surfaces soon degrade when exposed to wind, rain, dust, and snow. One way to circumvent these problems is to design a lightweight, flexible cloak that can act as a specialized reflector. Just such a cloaking system has been developed in the Tachi Labs at the University of Tokyo. A schematic overview of the Tachi cloaking system is shown in Fig. 3.3.

The system works by projecting on to the front of the cloak an exact rendition of the view that lies behind it. Accordingly, a digital camera is placed behind the cloak, and its output is directed to a projector aimed towards a flat, angled at a 45° transparent screen – this is the Pepper's Ghost part of the device. The flat screen allows the observer to see right through it, but it also reflects the projected background image through a 90° angle so that it illuminates the front of the cloak. The subterfuge is now complete, and the cloaking device is ready for use (Fig. 3.4).

The Tachi invisibility cloak is a neat example of applied technology – but at this stage it is only a demonstration device that works under restricted conditions. Indeed, the invisibility condition is only satisfied for an observer with their eye located within the small viewing cone set in front of the combiner (Fig. 3.3).

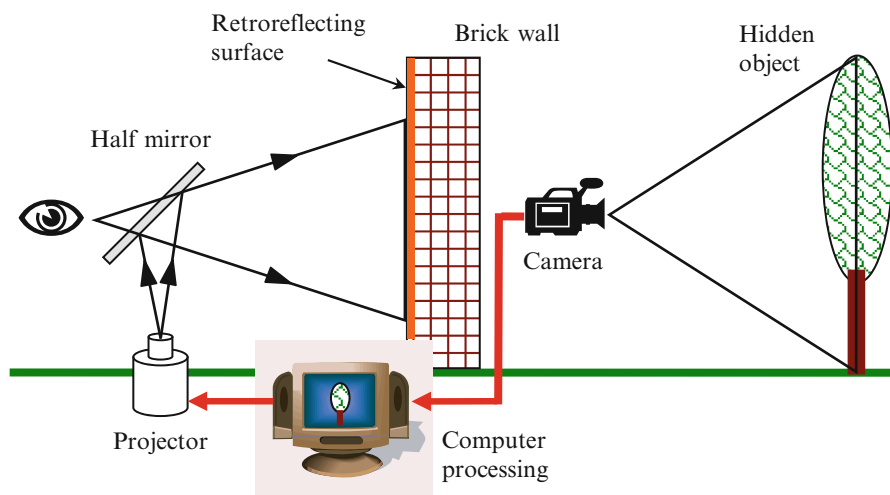


Fig. 3.3 Schematic layout of the front-projection Tachi invisibility cloak



Fig. 3.4 A real-life example of invisibility cloaking. The view occulted by the central figure is projected seamlessly onto the front of the cloak, making the wearer appear (partially) invisible (Image courtesy of Tachi Labs, The University of Tokyo)

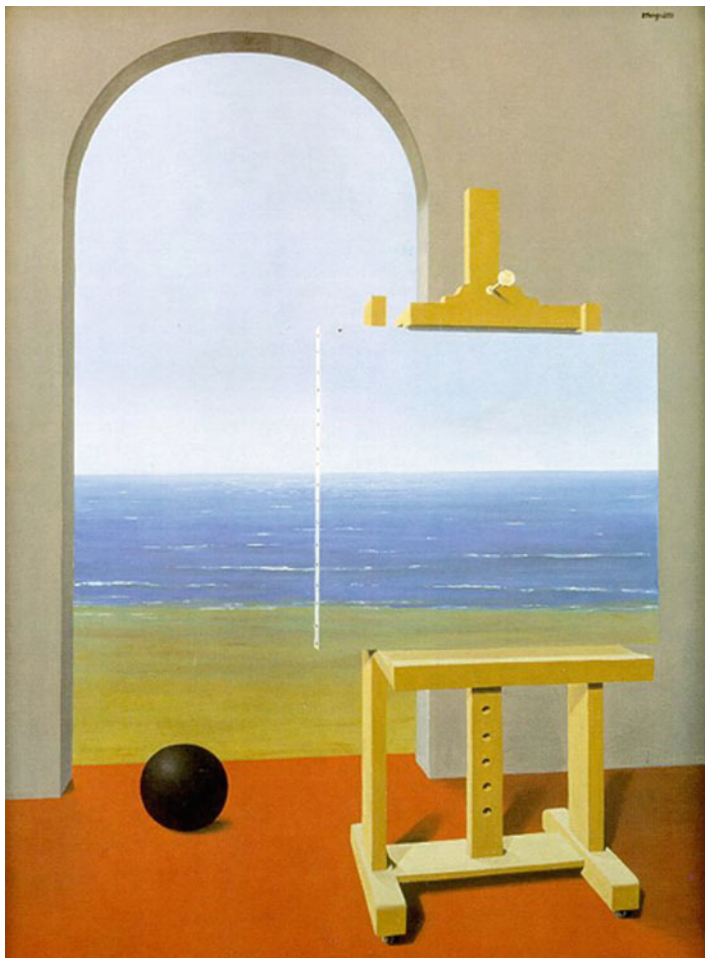


Fig. 3.5 “The Human Condition” by Rene Magritte. Painted in 1935, Magritte argued that it was actually irrelevant if the real scene behind the picture was different from the painting. The painted vista, however, must join harmoniously with the directly observed fraction of the outside scene or else the illusion is lost

That is, the observer needs to see light directly reflected back along the path of the projected image. As a practical, dynamic, and moving-observer invisibility device the Tachi cloak is rather limited, but it comes into its own when the circumstances are such that the observer and screen remain fixed – circumstances that painter Rene Magritte realized in his painting “The Human Condition” (Fig. 3.5).

A dynamic variant of Magritte’s illusion has recently been developed and tested by weapons’ manufacturer BEA Systems. In this case each side of a military tank is equipped with a series of cameras and display tiles made of light-emitting diodes. The image captured on each side of the tank is then transmitted to the optical arrays



Fig. 3.6 The adaptive E-ink camouflage developed by BEA Systems updates the image portrayed on each side of a tank, making it blend into the background (Image courtesy of BEA)

on the opposite side. In this way an external viewer is presented with the view that they would see if no tank were present – in essence seeing right through it (Fig. 3.6). The projected images need not be perfect, just good enough to fool a distant observer and make the tank a far more difficult target to engage on the battlefield.

A Moment's Reflection

In principle the cloak used in the projecting system could be made of any reflective material, but the Tokyo University research group used a material with a special reflective surface. Indeed, the finish was especially chosen for its ability to reflect light back along its original path. This behavior is somewhat different from the law of reflection as presented in Chap. 2, where the angle of incidence is equal to the angle of reflection.

This law still holds true, of course, but now we are more concerned with the reflection of many rays from an extended surface rather than a single ray from a single point. The manner in which light reflects off an object will depend upon how smooth or rough its surface texture is. A very smooth surface will generally produce a so-called specula reflection, while a rough surface produces a diffuse reflection. The differences between these two modes are illustrated in Fig. 3.7. In the case of the Tachi cloaking device neither specula nor diffuse reflection is ideal since most

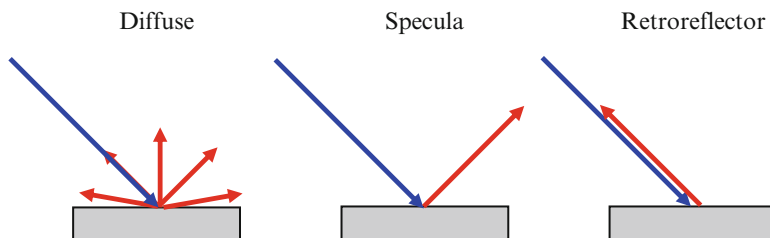


Fig. 3.7 Ray diagram illustrations of diffuse, specula and retroreflection

of the light will be reflected off the cloak at angles away from the observer. Rather, the surface required must be a retroreflector.

We have all seen examples of retroreflector paint surfaces but may not have realized it. Virtually all modern road signs, road median markings, and portions of the clothing worn by cyclists and emergency response personnel have retroreflector surfaces. Indeed, under any circumstances where high visibility is an asset, retroreflection surfaces are a definite help. They are also used in many laser ranging projects, and several of the Apollo lunar landing missions conducted during the late 1960 and 1970s deployed such devices so that highly accurate measurements of the Moon's motion could be made.

It is perhaps ironic, in some sense, that an invisibility cloak must have a surface that produces the brightest possible reflection and specifically a reflection that is almost exactly back along the direction of the incoming light ray. The retroreflection effect is achieved by coating a reflective surface with myriad small glass or plastic spheres (Fig. 3.8).

Never missing a chance of producing something beautiful, nature has also invented a form of retroreflecting surface, although from a human perspective one has to rise with the very larks, in the early morning, to see it. The surface is that of a dew-covered lawn of grass. Indeed, via the action of surface tension and surface minimization, the dew drops on short blades of grass will take on an approximately spherical shape, and these will act as approximate retroreflectors. The best time to see the effect of this watery retroreflection is just after sunrise. At this time, with the Sun low in the sky, the *heilighenschein*, or “holy glory,” can be seen. Standing with your back to the Sun look at your shadow and you will find that the retroreflecting dewdrops have produced an aureole of diffuse light around your shadow – you will literally appear to be glowing.

Cats' Eyes

We have all probably seen cats' eyes – both those made by nature and those fabricated by human hand. Indeed, it was by seeing his car headlights strongly reflected through the eyes of a nocturnally exploring feline (Fig. 3.9) that British inventor

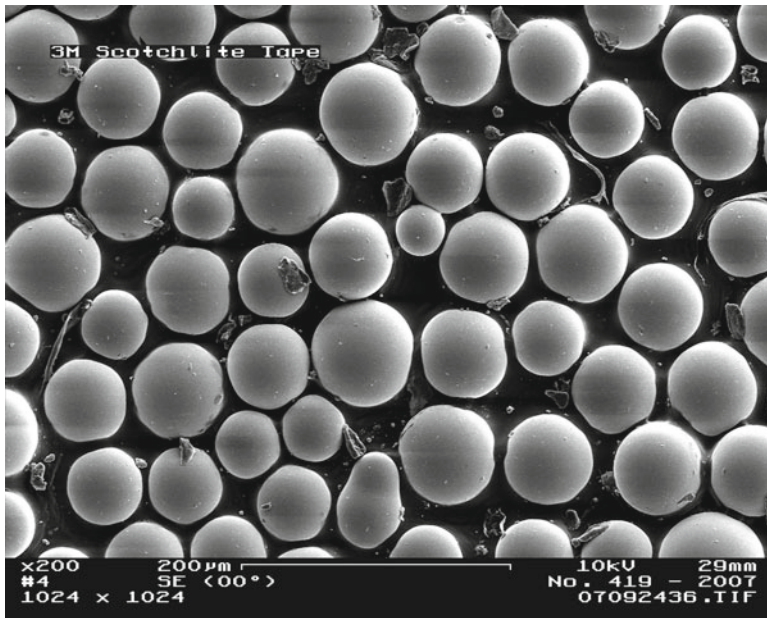


Fig. 3.8 A 100× magnification, scanning electron microscope, image of the surface of 3 M Scotchlite retroreflector tape. It is the small spheres, which are about 50 μm across, which give the material its special reflective properties (Image courtesy of Wikipedia commons)



Fig. 3.9 The cat's eyes reflection effect here demonstrated through the canine eyes of Mr. Vincent Noir

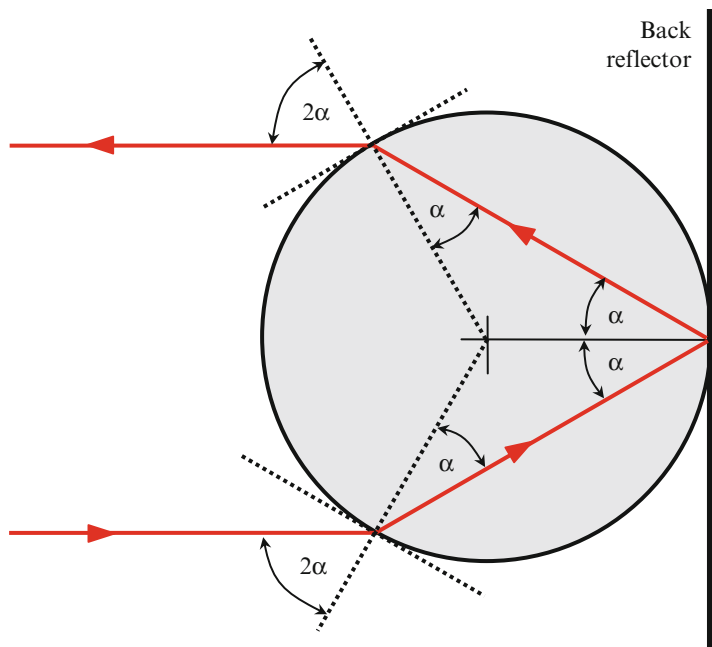


Fig. 3.10 Ray diagram illustrating the light path through a retroreflecting sphere

Percy Shaw came up with the simple idea that has saved countless numbers of human lives. The patent for his cat's eye safety device was filed in 1934, and now, by their many millions, they decorate the central medians and traffic lanes upon roads throughout the world.

The idea behind the cat's eye is to direct back to the driver a bright reflection of the headlamps. In this manner the eyes act as a visual cue to keep the driver in the correct road lane.

Shaw's cat's eye is essentially a single, centimeter-sized retroreflector sphere. The way the cat's eye reflector and retroreflector surfaces in general work is that most, ideally all, of the light that strikes them is reflected back along a path that runs parallel to the original, incoming light ray. In order for this to come about the light ray must be refracted twice and reflected once by the sphere. The geometry of a cat's eye lens is shown in Fig. 3.10.

At the location A where the light first strikes the sphere, the condition for the exit ray B to be parallel to the incident ray is that the angle of refraction is exactly twice the angle of incidence: $i = 2r$. From Snell's law the angle of refraction is related to the refractive index of the air (n_{air} , for the incident ray) and the sphere (n_{sphere} , for the refractive ray), with $n_{\text{air}} \sin(i) = n_{\text{sphere}} \sin(r)$. Now, since the refractive index of air is essentially one, so, the parallel condition between the incident and exit ray is dependent upon the refractive index of the sphere having a value greater than 1 but less than 2. It turns out that no constant refractive index sphere will perform as a true

retroreflector for all angles of incidence, but a detailed numerical analysis of the incident versus (near) parallel return ray characteristics of spheres reveals that the best fixed-value for the refractive index is $n_{\text{sphere}} = 1.85$.

In terms of constructing cat's eye lenses with material having a constant refractive index our discussion is now complete, but we shall come back to consider the possibility of constructing a perfect cat's eye lens later in this chapter.³

Seeing Through a Door

Question: When is a door not a door?

Answer: When it has a retroreflecting inside surface.

Augmented-reality technology (ART) first appeared in the early 1960s, and it developed in concert with the then new electronic computer industry. ART is all about bringing computer derived and analyzed data into the visual realm – augmenting, therefore, the perceived world of the observer. One of the key steps in enabling ART to enter the everyday realm was the invention of the head-mounted display (HMD) unit by pioneering computer scientist Ivan Sutherland in 1968. HMD devices have found extensive use in military operations as well as in the computer gaming world and are composed of two electrical display screens positioned in front of an observer's eyes – in most applications the observer sees nothing else but the screens. With HMD devices video images recorded at, say, infrared wavelengths can be displayed to an observer, augmenting, therefore, their visual perception.

It will be through the application of ART and HMD that the Tachi invisibility cloak will most likely be brought into our everyday world. The researchers at the Tachi Labs have already proposed and begun researching several commercial applications of their invisibility cloaking technology using HMD. The trick is to find applications where the observer's position does not change much relative to his or her surroundings – such as in a car or in an aircraft cockpit. Clearly the car or aircraft will be moving, but the view of the interior floor and paneling remains fixed to the inside. Figure 3.11 shows one application of this technology where the side paneling of a car is made to appear invisible, allowing the driver to have a clearer view of a bicycle – a particularly useful asset if driving along narrow streets. The manner in which the so-called transparent cockpit system works is illustrated in Fig. 3.12.

The idea illustrated in Fig. 3.12 is that the image projector is mounted to the observer's head, and that a computer continuously monitors the direction of the viewer's gaze. In this manner a correctly scaled image of the outside world can be projected, via the HMP, onto the retroreflector-covered interior wall of the cockpit.

³ The reader is challenged to consider the following thought experiment, or kitchen sink trial: *Do safety retroreflection jackets work effectively under water?* – Answer: no since the refractive index of water is greater than that of air, and this changes the refraction rule at the water – retroreflector interface.



Fig. 3.11 By being able to *see* through the car door a driver can more accurately gauge the distance to a cyclist or pedestrian (Image courtesy of the Tachi Labs, the University of Tokyo)

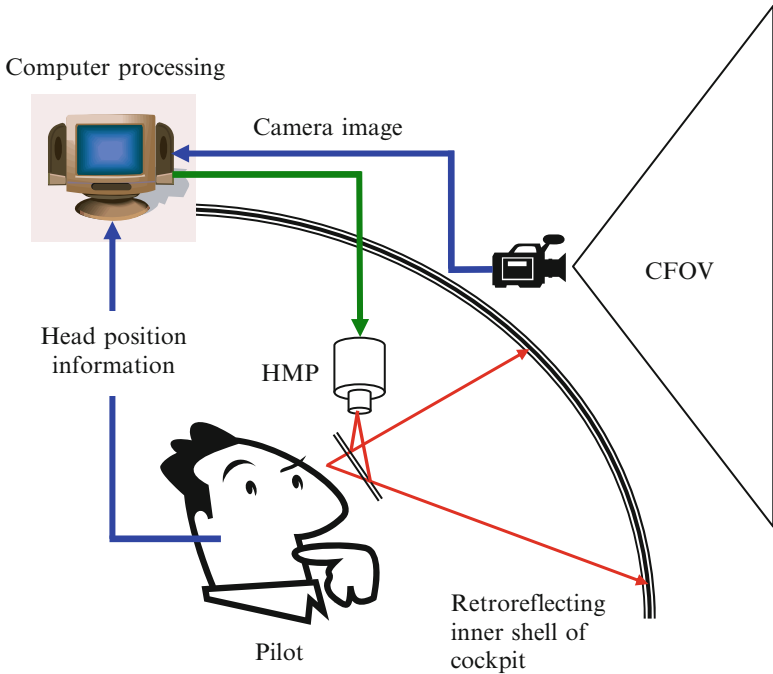


Fig. 3.12 A schematic diagram for the layout of the transparent cockpit system. The computer senses the head position and relays to the head-mounted projector (HMP) the appropriate scaled image to be projected onto the retroreflectors' surroundings. CFOV is the camera field of view

Writing in a recent review paper Takumi Yoshida (University of Tokyo, and the Tachi Labs) and his coworkers have commented that, “The ultimate goal of the transparent cockpit is that the driver should not feel as if *I am operating a vehicle*, but *I am running or flying by myself*.” To this Yoshida adds that this goal should be achievable within a few decades.

Alpha Blending to Half Transparency

There are times when being only half-hidden is useful, as the researchers at the Human Interface Technology (HIT) Lab at the University of Canterbury in New Zealand recently found out. A growing number of physical tasks in every day life, industry, and space exploration are being mediated by humans using computers – a process called telepresence. Under these circumstances the human might be completely removed from the mechanical system, as, for example, the drivers who control the Mars Exploration rovers *Spirit* and *Opportunity* – the distance between the vehicle and driver being in excess of 100 million kilometers on occasion. In contrast to this, other operators might be required to interact manually with the system – as in the case of a surgeon performing a minimally invasive operation. Here the surgeon uses images obtained with fiber optical links to see what they are doing inside of the patient, rather than opening them up for direct visual inspection. Typically under these latter circumstances the system is designed to combine two images – a background image (say of the patient’s cancer) and a foreground image (say of the surgeon’s tools). The question that then arises is what is the best way to combine the two images in order to help the surgeon?

A foreground object’s level of transparency, compared to that of a background image, can be described according to a numerical parameter called, let us say, alpha. An alpha of zero corresponds to full transparency (that is, the foreground object is completely invisible to the observer). At the other end of the scale, when the foreground object is fully opaque, the alpha term attains a maximum value of one. The level of preferred transparency has been a topic studied at the HIT Lab at the University of Canterbury in a series of studies in which a human operator is asked to perform a number of manual tasks while wearing a head-mounted display (HMD). In this test, the operator’s hand corresponded to the foreground object and the tasks (writing a name, moving an object, plugging in a USB memory stick, etc.) were to be performed within a fixed-background environment. The researchers varied the value of alpha, blending it to produce differing levels of transparency of the operator’s hand. Among all of the test subjects the preferred alpha levels were found to fall between 0.6 and 0.8. For values of alpha smaller than 0.5 most subjects felt that the image no longer corresponded with what their hand was doing – that is, what was seen on the image screen and what was felt by their hand seemed all ‘wrong.’ In addition, if the alpha value was set too low then the experimenters lost any sense of depth perception within the images shown in their HMD.

Crystal Vision and the Invisible Touch

An alternative way of combining image information is to use retroreflection projection technology (RPT). In this case, rather than blend images according to a specific alpha value, the data to be superimposed upon the background (a specific retroreflecting surface) is actually projected by a head-mounted unit. The basic design for an RPT system is shown in Fig. 3.13.

The Tachi Labs at the University of Tokyo have worked on a number of RPT systems and have even developed a specific head-worn projection unit, the X'tal visor, that uses a pinhole mirror to project the image of a virtual object onto a retroreflector screen (see Fig. 3.14). While RPT systems have many conceivable applications within industry and medicine, they can also be used to completely change the appearance of an object. In this latter situation we begin to enter the realm of optical camouflage, where the idea is to disguise an object by making it look like

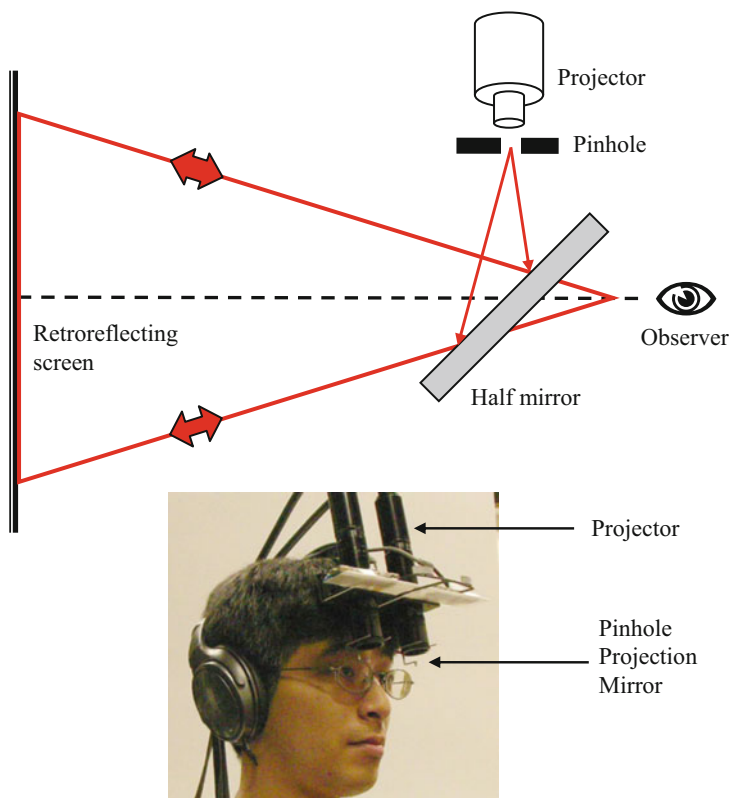


Fig. 3.13 Schematic design layout for an RPT system. A pinhole projection system is used to form an image on a retroreflecting surface via a Pepper's Ghost 45° angled plate of glass. The observer will see a combined view of the background and the virtual image projected onto the retroreflecting surface. The lower inset shows the prototype X'tal vision projection



Fig. 3.14 The X'tal (crystal) vision and RPT. In this case an optical image of an X-ray chart is directed towards a retroreflecting shirt, thus enabling the observer to see both a real and a virtual image at the same time

something else. In essence this form of invisibility is a Trojan horse, with the object being hidden from view in plain sight.

One of the main goals of ART is to enable an observer to interact with real-world objects in an enhanced fashion, and accordingly there are now a number of research groups actively exploring ways to extend the augmented reality environment to include the sense of touch. Such visuo-haptic displays, the word haptic being derived from the Greek for *sense of touch*, will allow, for example, a person to not only see virtual objects but to also touch and feel them. The Tachi Labs, as ever in the thick of things, have begun designing a visuo-haptic system in which a virtual object is projected onto a hand-held retroreflecting sphere. This system takes RPT into another realm and requires that eye, head and hand tracking must be taken into account so that the correct orientation of the virtual object is achieved. By projecting the image onto a retroreflecting sphere the observer's hands correctly occlude the visual display enhancing both the three-dimensional feel and sustaining a correct sense of depth perception.

The ultimate goal of the visuo-haptic RPT systems is to include shape approximation technology. Such devices would then enable the observer to feel *real* corners, bumps and surface textures. One device that the Tachi Labs have recently tested is that of the *Pen de Touch*. This hand-held, pen-like gizmo enables an observer, using an HMD, to explore and feel a virtual environment. The pen not only monitors the

observer's hand motions, but three motors inside of the pen provide the observer with a tactile sense – revealing hard surfaces, textures and edges. This very same technology may soon be applied to the ubiquitous cell phone, enabling a user to send tactile messages as well as visual and textual ones. Indeed, Shwetak Patel (University of Washington in Seattle) and co-workers have recently developed the SqueezeBlock – a mobile haptic device that can convey information by changing its stiffness. The device, for example, can be made to feel floppy when its battery is low or rigged when the battery is fully charged. Fabian Hemmert (Deutsche Telekom Laboratories, Berlin) has also experimented with designs that could enable a cell phone to change its profile as well as its center of mass – all tactile properties that the user can sense and accordingly respond to.

Augmented Reality

The cell phone, global positioning system, CCD camera, lap-top computer and web browser are all devices that shape our reality, and we increasingly explore the world and communicate through an ever-growing electronic network. Such devices not only enable us to experience the world, they also enable us to see it and explore it in entirely new ways – ways in which the body alone never could.

The Tachi Labs have already developed and demonstrated the ARScope (Augmented Reality Scope). This hand-held device enables the user to seamlessly see the world at different magnifications and in different colors. Shaped like a magnifying glass the ARScope uses cameras (on the back face), a retroreflection surface (on the front face) and a head-mounted projector system to display the image. While a prototype device has been constructed and displayed, future developments could include magnification capabilities, viewing at wavelengths of light outside of the optical (augmenting human eyesight) as well as active contact lenses (Fig. 3.15). Indeed, the applications are near limitless, or perhaps we could say that the applications are only limited by our imaginations.

While not an actual augmented reality device, the Federal Institute of Technology in Lausanne, Switzerland, has recently developed and begun to market a remarkable contact lens – called Triggerfish. In this case the lens actively monitors the condition of the wearer's eyes, transmitting the information wirelessly to a portable recording device. The contact lenses, in fact, measure the curvature of the cornea in each eye and the data transmitted back can be used to locate regions of internal pressure variation (such as those associated with glaucoma). A similar such smart contact lens has been developed by Babak Parviz at the University of Washington in Seattle, in this case, however, to monitor glucose levels in diabetes patients. With this new technology our eyes are not only the windows to the invisible soul within, they are also the transmitters of our physical health without.

Augmented reality devices remind us of the fact that sometimes what we think we see is not necessarily real – as Magritte long-age warned us (recall Fig. 1). And, while electronic images may not be true representations of our surroundings, we

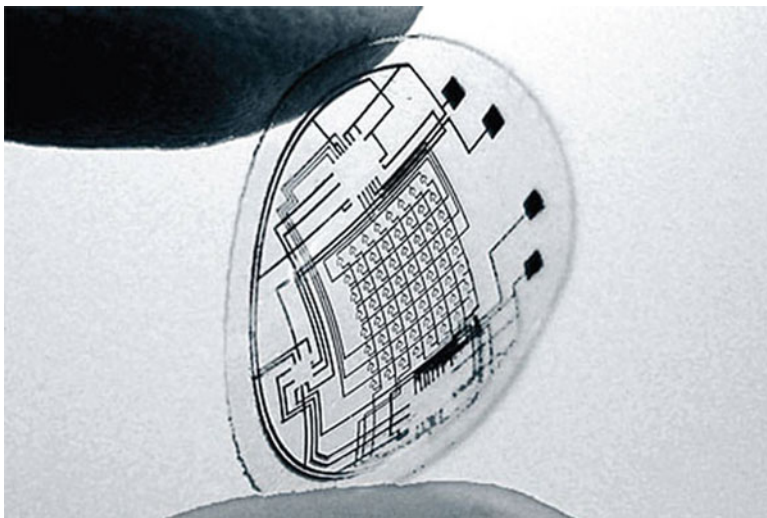


Fig. 3.15 Prototype augmented reality contact lenses developed by Babak Amir Parviz (University of Washington)

should also remember that nature can play tricks of the light upon us, producing through atmospheric channeling non-existent impossible vistas that for all of their apparent reality we may never hope to physically explore. These are the crepuscular worlds of mirages and *Fata Morgana*.

Fata Morgana

In modern usage the term *Fata Morgana* (Fig. 3.16) is reserved for those mirages produced by temperature inversions that run along an atmospheric duct (Fig. 3.16). Such atmospheric channels are responsible for many odd optical phenomena, but the term simply refers to an extensive region of the lower atmosphere in which light rays can become trapped and ducted. In order to form, such optical tunnels need a stable atmosphere and the trapping of a cold air mass close to Earth's surface by a warmer overlying region of air. Since the refractive index of air varies according to its temperature, so the light rays of distant, over the horizon objects can form refracted arcs directed towards the observer (Fig. 3.17). The back-projection of the ducted light rays will then appear to the observer as images located above the horizon.

The nautical tale of the Flying Dutchman, a legendary ghost ship supposedly doomed to sail the seven seas for all eternity, never being able to make a safe anchorage, may well take its origins from the unsuspected appearance of *Fata Morgana* mirages – reminding us that what we see, and the manner in which a given vista is interpreted, is very much in the eye of the beholder. Echoing this very point, Olaf

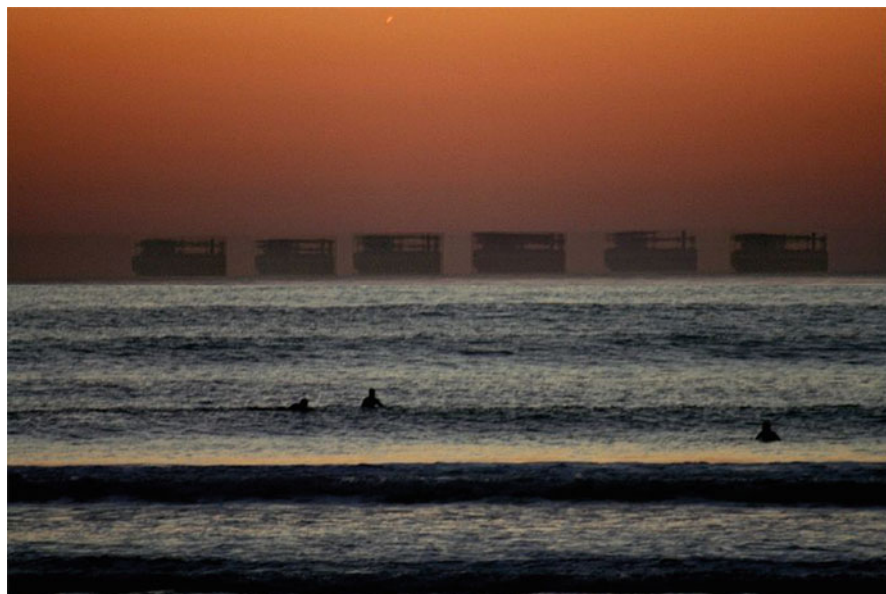


Fig. 3.16 A series of six photographs have been combined in this image to show the time variation in the Fata Morgana distortion of a distant ship (Image by Mila Zinkova)

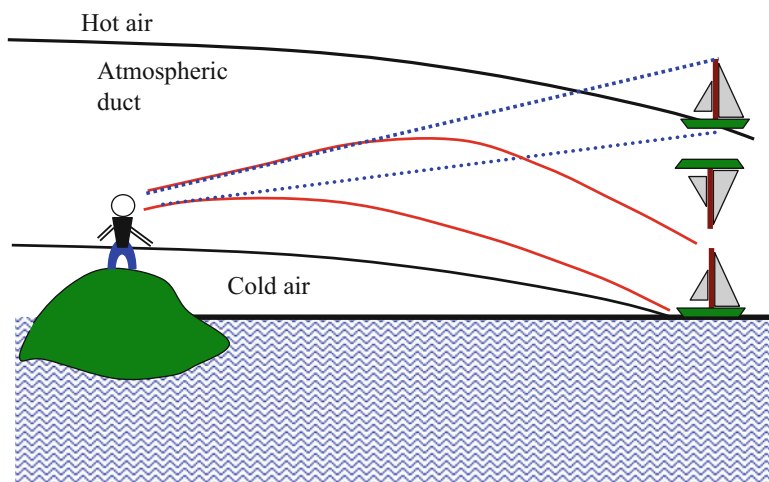


Fig. 3.17 Schematic ray diagram for the production of a Fata Morgana in an atmospheric duct supporting a temperature inversion layer. Only a few optical ray paths have been shown for clarity

Stapledon, in his remarkable flight of literary fancy, *Star Maker*, reminds us that, “imagination completed what mere sight could not achieve.” Not only, however, can ancient mariners be fooled by what they occasionally see, so indeed, are we all deceived, every day, by the setting Sun – as we shall explain next.

Earth's Atmosphere and the Sunset Deception

The idea that Earth's atmosphere can act like a giant lens was realized long ago by the ancients, but how to correct for its distorting effects is a far from trivial problem. The correction, however, it turns out, is always positive, making objects appear higher above the horizon than they truly are (see Appendix 2). The only sky location where there is no atmospheric correction in fact is that point directly above an observer's head – the so-called zenith location. Astronomers also tend to use zenith angle Z , the angle away from the local vertical, rather than the altitude A , the angle above the horizon, to describe the position of a star in the sky. This situation is illustrated in Fig. 3.18 where the atmosphere is approximated as a series of horizontal slabs, each of which has a slightly different refractive index.

No atmospheric refraction model is fully accurate when a star is close to the horizon (zenith angles close to 90°), but the average value is generally taken to be about 0.5° (as shown in Appendix 2). In other words, since refraction always makes an object appear higher in the sky than it actually is, a star that appears to be right on the observer's horizon is actually about 0.5° below the horizon. In other words the atmosphere enables us to see objects that are geometrically unobservable.

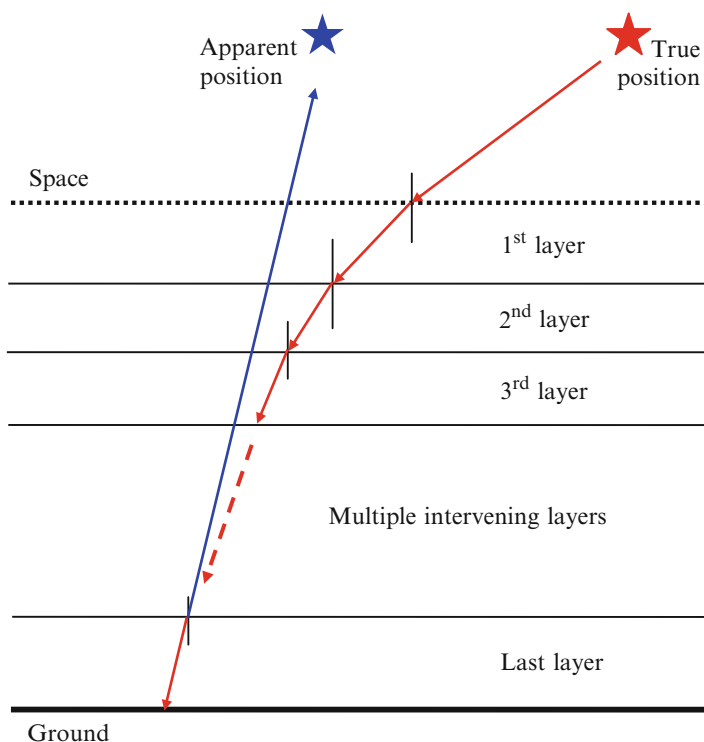


Fig. 3.18 The passage of starlight through a plain parallel, multi-layered atmosphere

By pure coincidence it transpires that the Moon is about 200 times smaller than the Sun, but the Sun is 200 times further away from Earth than the Moon. Observationally, what this means is that the angular diameter of the Sun and the Moon are the same and about 0.5° for terrestrial observers.⁴ Remarkably, the entirely unrelated coincidences concerning the horizon refraction limit and the angular size of the Sun and Moon transpire to produce a setting illusion. When we see either the full Moon or Sun setting, with their disks just touching the horizon, they have in fact already set (geometrically speaking) below the horizon. In this sense, the atmosphere makes the invisible visible.

The Green Flash

Sounding remarkably like a comic book superhero, the green flash is actually an atmospheric phenomenon that, in principle, can be seen on any day of the year, requiring just a clear view of the horizon and a visible sunset. As the name suggests it is a reasonably rapid phenomenon and one that is over in just a few seconds – the green flash is literally the last vestiges of the Sun's disk visible to an observer before it finally blinks-out of view. It lasts just long enough to leave one in doubt as to whether it was actually seen.

The appearance of the green flash comes about partly as a result of Earth's atmosphere acting like a lens, bending light waves around its curved surface, and partly because the atmosphere acts just like a prism, separating light into the colors of the rainbow. This latter effect comes about because the refractive index of light is not only sensitive to atmospheric temperature, it also varies with wavelength. It turns out, as a consequence of this behavior, that light rays corresponding to the wavelength of blue light curve more in the atmosphere, for a given temperature, than those corresponding to red light. What this means is that the red (longer wavelength) waves are obscured from view (that is, fall below the observer's horizon) before the blue ones. Blue light, however, is readily scattered by the multitude of dust grains and molecules in Earth's atmosphere, via a process known as Rayleigh scattering.⁵ What is technically the final blue flash of the setting Sun is actually invisible to us. Accordingly, it turns out, the most obvious final hue of the Sun, as it dips below the horizon, is that corresponding to light with a wavelength slightly longer than blue light, and this means that green light predominates.

⁴ The angular width of your index finger, when viewed at arm's length, corresponds to about 1° on the sky. The full Moon will therefore have an angular diameter of about half that of your index finger held out at arm's length. Check this result the next time that the full Moon is observable.

⁵ The daytime sky is blue because of this very same process. In contrast, on the planet Mars the daytime sky is a yellow-brown color, due to a predominance of suspended dust particles and the fact that its low density atmosphere provides fewer molecules for the Rayleigh scattering process to work upon.

How to Catch an Atom Bomb

Viewing the Sun directly should never be attempted without great care – its overwhelming brightness is literally damaging and painful to gaze upon. Not only this, the brightness of the Sun so overpowers the eye's receptors that nothing but a glaring miasma of light is experienced anyway.

The best times to view the Sun with the unaided eye, as every seaside holiday maker knows, are during sunrise and sunset, when its light is reddened and diminished. Indeed, the sunset phenomenon and the Green Flash are only visible to the human eye because of the fact they are seen through an extensive column of atmospheric gas and dust (see Appendix 2), and because they manifest themselves on timescales varying from a few seconds up to several tens of minutes. In some situations, however, not only is an event far too bright to look upon, it can also unfold on a mercurial timescale, changing far too rapidly for the brain to follow. The U.S. military found itself in just such a situation in the mid- to late-1940s when developing the first atomic bombs. Not only could no human observer look upon and unravel the sequence of events, neither could any conventional camera system. The consequences of the detonation were easy enough to record and measure, but what about those first moments of atomic genesis?

The task of developing a high-speed shutter became paramount, but MIT professor of engineering Harold “Doc” Edgerton eventually solved the problem. Indeed, Edgerton is often described as the father of high-speed photography; he was a pioneer of the stroboscopic technique used to capture the now iconic image of a falling milk drop, and he helped develop the rapatronic (short for Rapid Action Electronics) camera. The ultra fast shutter speed of the rapatronic camera was achieved by using two polarizers (recall Chap. 2) situated either side of a Kerr cell. The Kerr effect was discovered by Scottish physicist John Kerr in 1875 and relates to the observation that the refractive index and polarization angle of some dielectrics and fluids can be altered by the application of an electrical field.

The shutter arrangement developed by Edgerton set the two polarizers at 90° to each other. In this normal mode, therefore, no light will pass through both of them and the shutter is effectively closed. By applying an electric field across the Kerr cell, however, the polarization angle of the light within the cell will be rotated and accordingly some will be able to exit the second polarizer – the shutter is open. Since an electric field can be switched on and off extremely rapidly, the time interval during which the shutter is open can be precisely controlled.

Figure 3.19 shows the awesome beauty behind the development of atomic bomb blast waves. Such images reveal otherwise invisible features – cavities and mottling form, it turns out, as a result of density inhomogeneities within the bomb casing, and ‘legs’ are sprouted through an interaction with the tower support cables – an effect that the researchers dubbed the “rope trick.” The brain struggles with the interpretation of the images shown in Fig. 3.19; the almost playful, amoeba-like appearance of the juvenile blast wave being contrasted against the deadly horror of the full-grown adult. Perhaps there are, after all, some things that we should never hope to see.

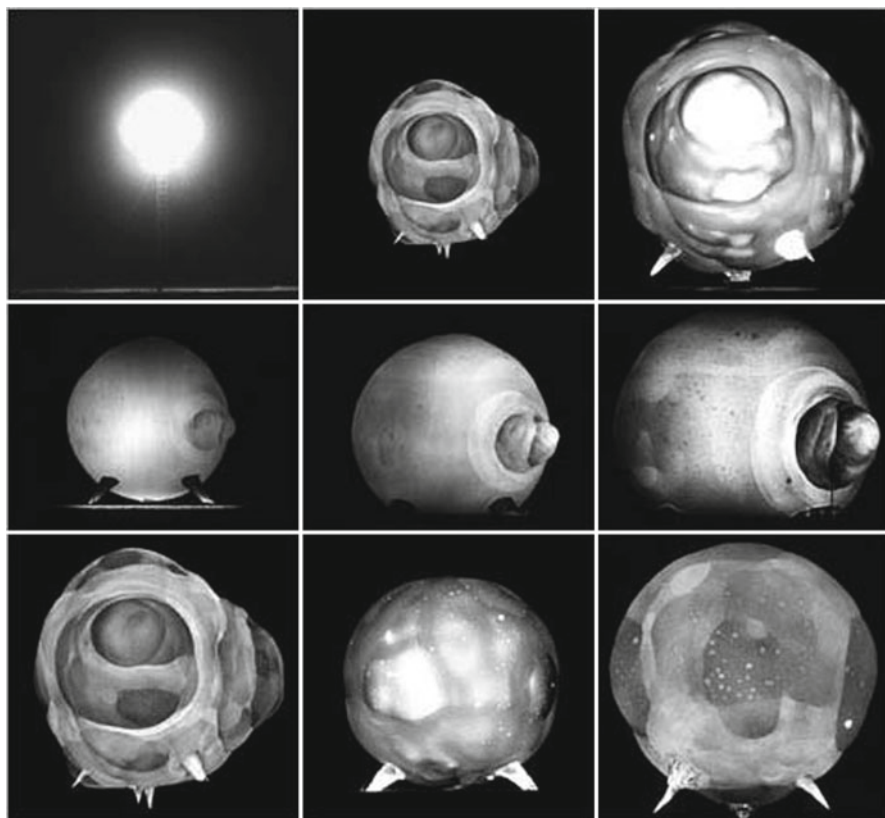


Fig. 3.19 The devastating beauty of atomic bomb blast waves is revealed in this series of high-speed photographs. The images show events unfolding within the first few thousands of a second after detonation, when the blast wave is smaller than a few hundred meters across

Mirages: Inferior and Superior

In contrast to the rapid Green Flash and the atomic bomb blast wave, the tremulous mirage is perhaps one of nature's most entertaining and pleasing tricks of the light. They are, however, impossible vistas, produced by nothing more than hot air. They fool the brain and they deceive the eye. They are the grandest of optical illusions writ large upon the landscape.

The refractive index of air, as we saw in Chap. 2, is generally taken to be one, but in reality it varies in a complex manner according to atmospheric temperature and density. The atmospheric density is at its highest at ground level and accordingly the speed of light will be at its slowest there since this is the region where the refractive index is at its highest. While the density in general decreases with atmospheric height it is in fact dependent upon the variation in both the atmospheric temperature and pressure. How these two terms vary will depend upon local weather and ground conditions, and

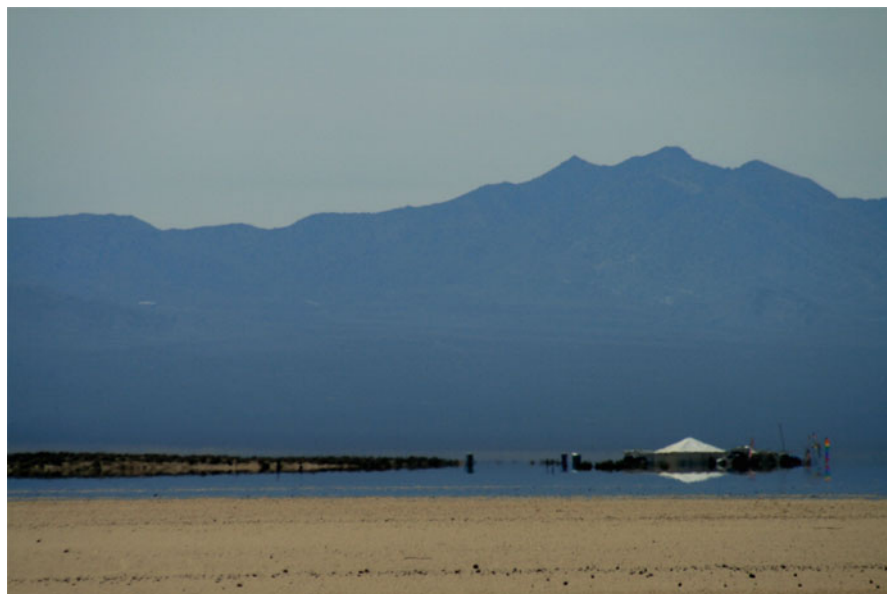


Fig. 3.20 An inferior mirage of a tent and the sky (Photographed in Primm, Nevada, by Mila Zinkova)

sometimes, when the temperature close to the ground is excessively hot or very cold, strange optical phenomena such as the Fata Morgana (recall Fig. 3.16) can be seen.

In broad brush form we can describe the atmospheric refraction effect on light as giving it the tendency to bend towards regions of higher density. Since the atmosphere is a near perfect gas it turns out that for a fixed pressure so the density and temperature vary in step with each other but in the opposite fashion – if the temperature increases so the density must decrease and vice versa. Accordingly, it is how the temperature varies with height close to the ground that determines the type of optical mirage that can be seen. If the ground temperature is much warmer than the air above it then a so-called inferior mirage can form (Fig. 3.20). If, on the other hand, a layer of warm air settles over a cold surface (such as on a lake or the sea) then a superior mirage can be seen (Fig. 3.21). The terms inferior and superior are not applied according to any specific image quality, but rather to whether the observer sees the image below (inferior) or above (superior) the miraged object.

That a mirage is produced by a region of smoothly varying refractive index tells us something useful about the way in which light rays can be controlled. Indeed, the lesson we learn from the formation of mirages is that the direction of a light ray can be changed gradually, in an unbroken fashion. This is in contrast to reflection and refraction, where the light ray remains perfectly straight but abruptly changes its angle of direction at the mirror interface, or refraction boundary. This gradual change in the refractive index is something that lens makers have long tried to emulate, but the fabrication processes are highly complex.



Fig. 3.21 A superior mirage of boats entering Victoria Harbor, British Columbia, Canada (Photograph by Craig Celments)

Robert Hooke discussed the idea of (in the modern vernacular) gradient index (GRIN) lenses in his 1665 *Micrographia* (recall Chap. 2), and he even suggested ways in which they might be constructed layer by layer. As with many of his projects and ideas, however, Hooke never managed to follow through with an actual result. Even if he had tried to fabricate a GRIN lens, however, the quality of glass and technology available in the mid-seventeenth century would probably not have produced any remotely acceptable end product. Indeed, the first GRIN lenses were not to be produced until the first decade of the twentieth century, when Robert Wood⁶ used a candle dipping method to produce a cylinder-shaped lens out of a liquid gelatin solution. Even to this day no true, optical wavelength, spherical GRIN lens has been manufactured. In spite of this technological short fall physicists and mathematicians have been fascinated by the special properties of GRIN lenses (as we shall see below).

A Cosmic Mirage

While the direction of a light ray can be changed abruptly via reflection and refraction, and smoothly in a GRIN medium, it can also be changed while simply propagating through empty cosmic space. This latter directive, however, is not brought

⁶ We met Wood earlier, in Chap. 1, as the first person to publish a research paper on the invisibility of transparent objects.



Fig. 3.22 A gravitational ring (SDSS J120540.43+491029.3) imaged by the Hubble Space Telescope (Image courtesy of NASA/ESA)

about by a variation in the refractive index (recall, $n_{\text{vacuum}} = 1.0$ by definition) but by the very curvature of space itself – this is the remarkable domain of gravitational lenses. Indeed, sometimes the human imagination can fly far beyond its bodily confines, travel deep into the primordial depths of the universe and see something before the human eye ever does. The young Albert Einstein had just such an imagination, and he saw in his mind’s-eye the possibility of a gravitationally lensed ring of light decades before any astronomer was able to photograph one (Fig. 3.22). Once again, for all of the remoteness of space, we are dealing with a mirage. These mirages, however, are now writ large upon the cosmic landscape.

The idea of gravitational lensing builds upon Einstein’s theory of general relativity. This theory describes how massive objects literally change the structure (that is, the geometry) of space (recall Chap. 2), and since light is constrained to move according to the geometry of space, so curved space means curved light rays and this enables the possibility of cosmic mirages. Astronomers detected the first gravitational lens in 1979, but the first Einstein ring was not discovered until 1987.

Einstein was clearly fascinated by the idea of light being directed by gravity and the possibility of gravitational lensing. He first considered the idea in 1912, developed it further in 1916 (in his initial general relativity paper) and then came back to

the concept in 1936. In essence a gravitationally lensed object is a mirage, since we only see it because the light rays are brought into our line of sight by an intervening massive object. If the massive lensing object, were not there then our view would be entirely different.

The appearance of the lensed object will vary according to its angular separation in our line of sight from the lensing mass, along with the mass and mass distribution in the lensing object. An Einstein ring requires that the lensed and lensing objects are exactly in our line of sight and the mass distribution within the lensing mass is smooth and spherical (Fig. 3.23).

An Einstein ring is both a mirage and an anamorphic print on the sky (recall Fig. 1.9). Incredibly, by analyzing the shape, size, color and brightness of the ring not only can astronomers recreate the shape of the background galaxy, unraveling the anamorph, but they can determine the mass of the lensing, foreground galaxy or, as is more often the case, cluster of galaxies.

In addition to measuring the mass of the lensing object, the physical distribution of the mass within the lensing object can also be determined. It is by these latter studies that astronomers have been able to map out the spatial distribution of dark matter⁷ (Fig. 3.24). There is a wonderful process at play here, and it beggars the imagination to think that the large-scale distribution of invisible dark matter can be traced out through the analysis of anamorphic images of miraged galaxies whose light rays have been directed towards Earth by gravity and the curvature of space.

Proving Einstein Right

The most famous cosmic mirage and gravitational lens in the history of science must surely be that observed during the total solar eclipse of May 29, 1919. Indeed, observations of this eclipse (Fig. 3.25) made from Sobral, in northern Brazil, and the Island of Principe, off the west coast of Africa, provided the first direct experimental confirmation of Einstein's general theory of relativity.

When he published his first paper on general relativity Einstein bolstered his argument by explaining, for the first time, the reasons behind the observed rotation rate shift in the orbit of Mercury – previously a long standing anomaly under standard Newtonian theory. He noted, however, that the Sun was massive enough to warp the space around it and that this should cause a measurable shift in the position of any stars that chanced to be close to its limb in our field of view. When the Sun is in full view, of course, the stars are invisible and the gravitational effect goes unseen. If observations are made at the time of a total solar eclipse, however, then the otherwise

⁷ It is not known what dark matter is made of and it is only recognized at the present through its gravitational influence on ordinary matter – specifically stars and galaxies. The name dark matter was chosen since it does not interact with or emit electromagnetic radiation. It is therefore invisible to our eyes and any telescope that we can currently make.

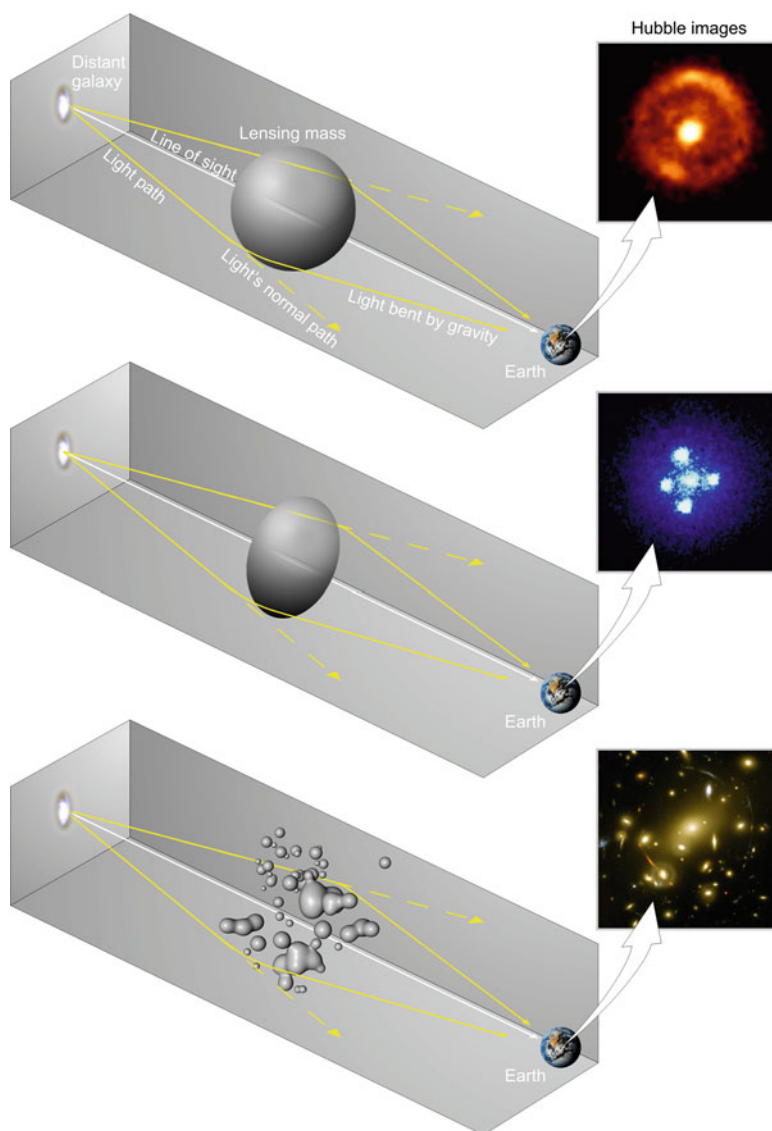


Fig. 3.23 How an Einstein ring is produced through gravitational lensing (*Top*). An ideal Einstein ring-forming gravitational lens (*Middle*). When the lensing mass is not spherical, multiple images rather than a ring are formed, and when there are many lensing objects (*Bottom*) the image of the distant object is stretched into long arcs (Image courtesy of NASA)

hidden stars can be viewed. By comparing two photographic images of the same star field, one taken during the eclipse and the other at nighttime, the star positions can be compared and any shift in the location where the Sun was during the eclipse can be measured.

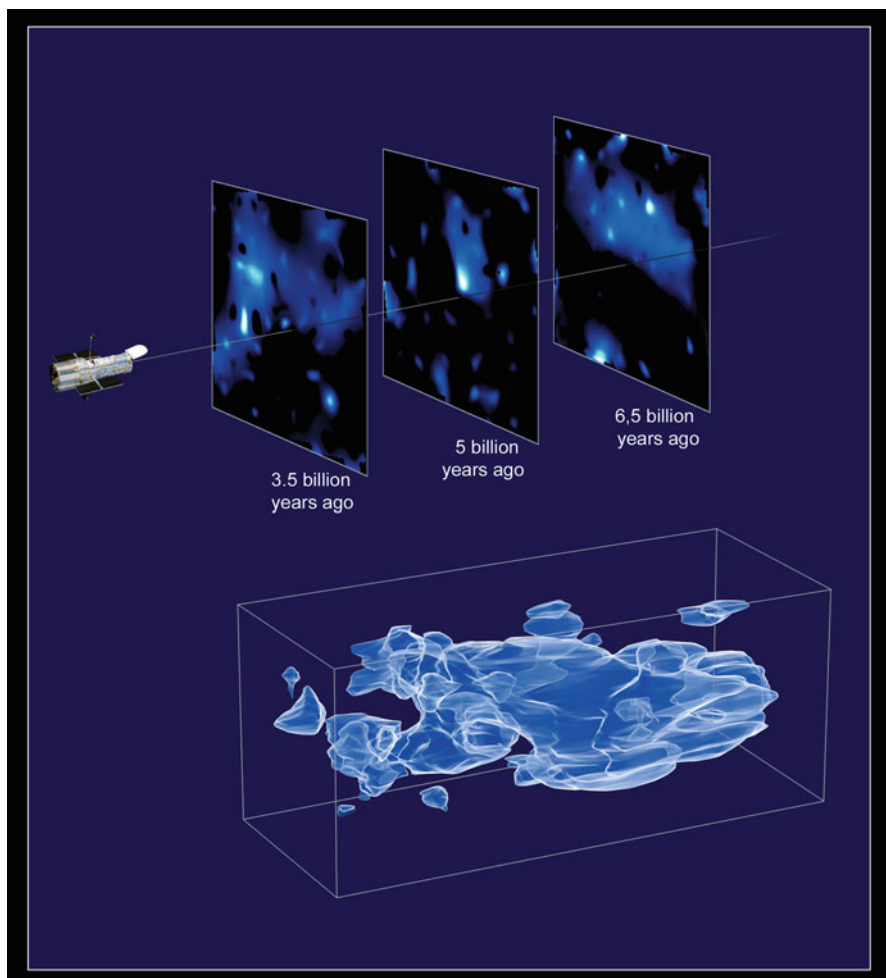


Fig. 3.24 The large-scale distribution of invisible dark matter, deduced by unraveling the gravitational distortions (anamorphic mirages) of background galaxies observed as part of the Cosmic Evolution Survey. The slices at the *top* show the 2-D sky concentrations of dark matter at three look-back times and distances (Image courtesy of NASA/ESA)

The results from the 1919 Eclipse Expedition were read to the assembled Fellows of the Royal Society in London on November 6, 1919, and the analysis paper was published in the January 1920 issue of the *Transaction of the Royal Society*. It was an astounding result, and within the paper the authors, F. W. Dyson (The Astronomer Royal), A. S. Eddington and C. Davidson, concluded that their measurements, “leave little doubt that a deflection of light takes place in the neighbourhood of the sun and that it is of the amount demanded by Einstein’s generalised theory of relativity, as attributed to the sun’s gravitational field.”

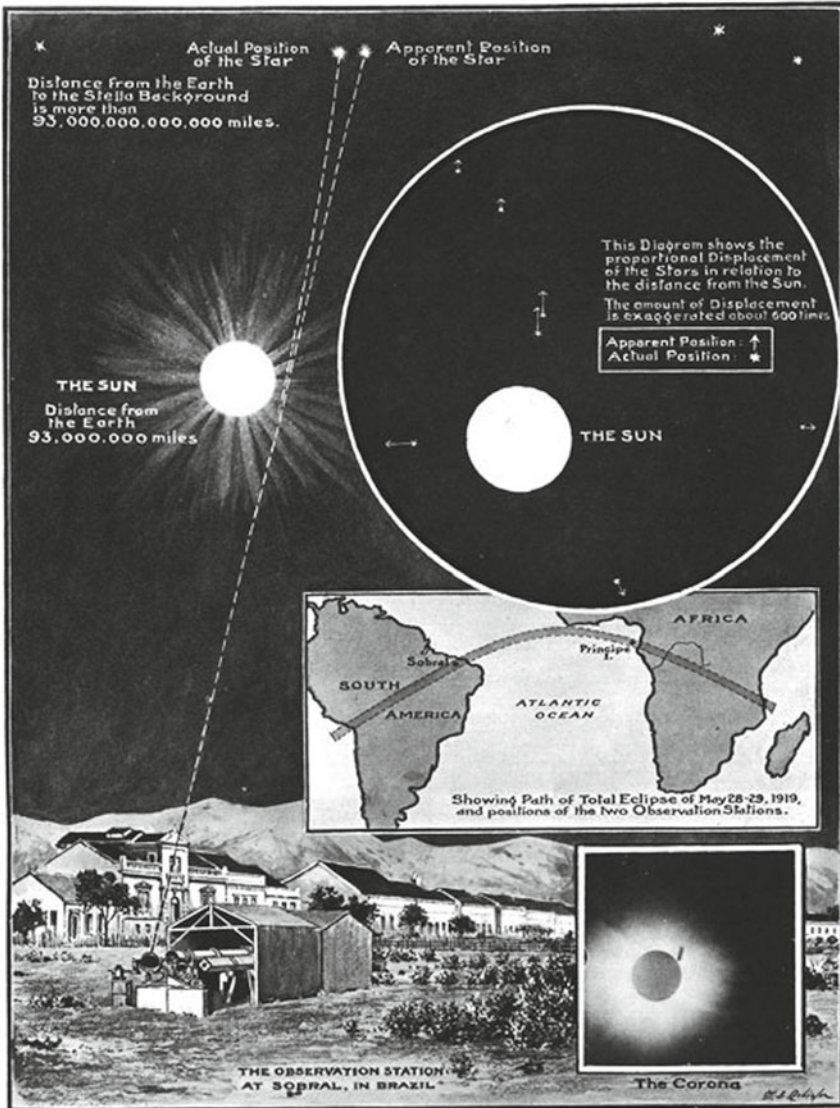


Fig. 3.25 A graphical explanation of the 1919 eclipse expedition to test Einstein's theory of general relativity as shown in the *Illustrated London News* for November 22, 1919

The Impossible Sunrise

Not only can light be steered by gravitational lenses it can also be constrained to move along invisible temperature-controlled corridors within Earth's atmosphere. Such atmospheric channeling is not common, but it can result in seeing an otherwise impossible sunrise.

The idea was simple enough. If the trade routes to the Orient around South Africa's Cape of Good Hope are being denied to you then, from Europe at least, head northward and search out a northeast passage over the top of Russia. Just as with the treacherous northwest passage over the top of Canada, many expeditions set out with the aim of finding navigable paths across the easternmost top of the world, through the Arctic sea. In the late sixteenth century Dutch polar explorer Captain Willem Barentz tried on three occasions to find a northeast passage. He failed, but he did see something that was supposedly impossible – a sunrise that shouldn't have been.

The journal of Gerrit de Veer, who traveled with Barentz, recorded on January 24, 1597 that, while trapped in an ice flow off the coast of the island of Novaya Zemlya, he and two shipmates briefly saw the Sun rise above the horizon. This was a surprise because the Arctic night was not due to end for a further 2 weeks – the Sun should not have been visible at all that day. This was not just a one-off display, however, and additional observations of impossible sunrises were noted by de Veer during his enforced stay in the Arctic ice. While initially questioned and even denied as being reliable, the veracity of de Veer's account was supported by other explorers working in both the Arctic and Antarctic regions. The phenomenon is now generally known as the Novaya Zemlya effect and is, in fact, the result of the formation of an atmospheric light duct, or waveguide, composed of a cool layer of air sandwiched between layers of warmer air (Fig. 3.26). The Novaya Zemlya effect is not just a mirage. It is a mirage of an otherwise invisible Sun located well below the observer's geometrical horizon.

The conditions required to produce the Novaya Zemla effect are only rarely satisfied, requiring the establishment of an extensive stable air mass and the formation of a thermocline – a region of warm air sitting on top of a layer of cold air. As illustrated in Fig. 3.26, these conditions allow the light rays from the Sun, situated well below the horizon, to wend a serpentine path along the waveguide, bringing them ultimately to a distant observer's gaze.

The light ducting in the Novaya Zemlya effect is one of nature's analog versions of the fiber optic cable – a long, thin strand of glass that can trap and transport light rays. The observation that light could be channeled along such strands of glass dates back to ancient times, but the first scientific investigations into the physics of the phenomenon were conducted in the mid-1840s. Fiber optic cables work according to the principle of total internal reflection. This phenomenon occurs when light is incident upon a boundary separating two transparent materials at a very shallow angle. Provided the light is traveling from a region of higher refractive index to a lower one it will undergo total internal reflection rather than pass through the boundary via refraction to the outer transparent zone. In an optical fiber a thin inner core of silica glass is wrapped by a cladding of lower refractive index glass such that light entering one end of the inner fiber will zigzag along its length, undergoing total internal reflection each time it intercepts the core-cladding boundary.

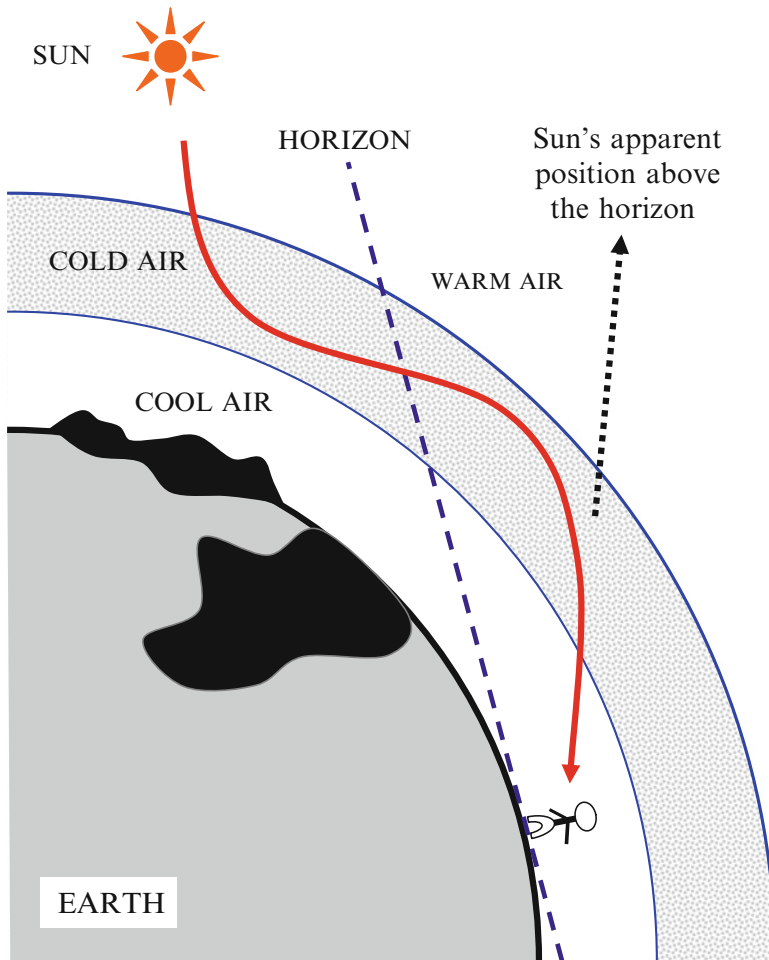


Fig. 3.26 Schematic diagram of the light ray path through a Novaya Zemlya-style atmospheric duct

Optical Fibers and Deep-Sea Sponges

Since the early 1970s countless thousands of kilometers of fiber optical cable have been laid over and under much of Earth's surface, and it is now through these great spider webs of crystal and plastic that we communicate. Whether by telephone, fax machine or computer, in the modern era, we literally talk to one another via the snake-like reflections of light bouncing down glass strands no more than a few millimeters in thickness. For all our human innovation, however, nature long ago discovered the utility of thin silica fibers. Indeed, the deep-sea glass sponge *Euplectella*, also known



Fig. 3.27 The Venus Flower-Basket, a deep-sea sponge called *Euplectella*. Each sponge is composed of silica fibers and is about 20 cm long (Image courtesy of NOAA)

as the Venus flower-basket (Fig. 3.27), is thoroughly remarkable for its ability to weave an intricate trellis-like cage from skeletal strands of silica fiber.⁸

In addition to growing its silica bordello, the *Euplectella* also fabricates an anchorage crown of numerous single strand silica spicules. Incredibly, these silica spicules are near-perfect optical fibers. Reporting on this discovery in the August 21, 2003, issue of *Nature*, Vikram Sundar (Bell Laboratories and Lucent Technologies) and co-workers noted that the spicules are composed of an inner core, just a few micrometers in diameter, with a refractive index of 1.46. This core is surrounded by a cladding, about 60 μm across, with a refractive index of about 1.43.

The spicules, it turns out, are a near perfect match to modern state-of-the-art optical fibers – produced, it must be added, by complicated sets of machines housed within giant, ultra-clean factories. Remarkably, the gentle deep-sea glass-sponge evolved a means of growing optical fibers millions of years ago. The research being conducted on the glass-sponge spicules at the Bell Laboratories is being pursued in the hope of learning how *Euplectella* manages to grow its fibers in the cold temperatures of the deep ocean. The researchers also hope to learn how the sponge fibers resist the crippling growth of internal fractures, the erstwhile bugbear and limiting feature of commercially produced optical fibers.

⁸ This cage, it turns out, is often co-opted to function as the love-nest for two mating shrimps – hence the sponge's more graphical common name.

From High Fashion to Cuttlefish and Squids

Light channeling via fiber optical cable is perhaps a weak form of optical cloaking in that the fibers physically manipulate the direction in which any interior light ray will travel. The fiber can make multiple loops and bends, and the light rays, trapped within the inner core, dutifully follow the path that has been laid down for them. In a limited field of view sense, an optical fiber could easily mimic the mirror cloaks described at the beginning of this chapter (recall Fig. 3.1). There is more, much more, however, that can be done with a set of optical fibers – the transmission of light and digital data is just the beginning of the story.

If clothes, at least according to author Mark Twain, make the man, then fiber optic fabrics will possibly make the wearer the advertising man of the moment. Indeed, fabrics containing plastic optical fibers have been manufactured for a number of years, and fashion designers have begun to incorporate them into their creations. Likewise advertising agencies have not been slow to pick up on the idea of using lightweight fabrics containing LEDs as well as optical strands, to make eye-catching and dynamic designs. Indeed, flexible screens made of woven optical fibers were developed and introduced by France Telecom R&D in 2002, and these wearable billboard screens can be used to display animated graphics and visual information.

In an example of science fact following science fiction, in October of 1996 a project initially sponsored by the U.S. Navy was initiated at Georgia Tech. The goal of the project was to fabricate a Smart Shirt – smart in this case meaning electronically wired. Under the guidance of Sundaresan Jayaraman, who is a professor at the school of Textile and Fiber Engineering, the aim was to produce a T-shirt woven from optical fibers and electrically conducting threads, with the intent that such a garment could monitor the health of the wearer (as a soldier, rescue worker or elderly patient) and in principle it could also send messages back to the wearer. From a military perspective it could also operate as an optical camouflage surface. The Smart Shirt is now licensed to and marketed by Sensatex in New York, who is developing a range of medical monitoring applications.

Writing in the August 2010 issue of the journal *Personal and Ubiquitous Computing*, Rehmi Post and Lit Whaal, at the MIT Media Lab, announced the next important step in active clothing design by describing the successful construction of a Teflon-nylon energy-harvesting fabric (called Sp4rk13 – which in real English reads as Sparkle) that uses the wearer's motion to generate electricity. This new fabric is essentially a wearable (shock proof) Leyden jar (to be described in Chap. 4), and the ultimate aim is to generate sufficient energy through it to power computing devices.

The new fiber optical fabrics and flexible screens potentially offer an alternative and perhaps more adaptable medium to the retroreflector fabrics presently used in optical camouflage and augmented reality devices; the main difference between the two media is whether the camouflage design is projected from an external device, onto the fabric surface, or generated internally, as it were, at the fabric's very surface. Time will tell how all these technology are going to play out, and while optical



Fig. 3.28 Two cuttlefish “cuddle” at the Georgia Aquarium (Image by David Iliff)

camouflage has clear military application it also has high survival selectivity – an attribute that the forces of evolution and natural selection learned about long ago (recall Fig. 2).

Perhaps the most spectacular creature to evolve an optical camouflage cloak and dynamical communication system is the humble cuttlefish (Fig. 3.28). While diminutive in size, the cuttlefish is a formidable hunter and a highly territorial fighter in the mating season – and it uses its dynamical camouflage to help with the former and announce its intentions in the latter. A cuttlefish is able to change its skin texture and color through the use of skin cells called chromatophors. These specialized pores can rapidly change their size and reflectance properties, enabling the cuttlefish to change its color, display complex patterns of stripes and splotches, and even flash brightly. All of these capabilities are used to confuse predators, bedazzle prey before delivering a wickedly fast *coup de grâce* and to signal the state of emotions to other cuttlefish. The cuttlefish is, in fact, the ultimate aquatic mood ring. The future use of similar such technology by humans using fabrics made of fiber optical cloth and integrated circuit materials is perhaps unclear, but soon we might literally be able to wear our (metaphorical) hearts on our sleeves – displaying, for better or for worse, our otherwise invisible emotions.

While technology inspired by cuttlefish chromatophors might eventually be used to broadcast the state of our minds, the leucophores of the Hawaiian bobtail squid might further be used to camouflage them. This particular cephalopod is a nocturnal hunter, and it uses both rapid coloration change and bioluminescence to communicate with its mates and to hide from predators. The skin of the Hawaiian bobtail

squid is a veritable mind-controlled camouflage cloak, and it was recently found that one of the key components in its biophysical defense system is a protein called reflectin. This particular protein forms part of the leucophore cell, and it is a remarkable broadband reflector.

Researchers at the Air Force Research Laboratory (AFRL) in Ohio have been actively synthesizing and investigating the properties of reflectin to see if its secrets can be learned and emulated. Ryan Kramer and colleagues at AFRL have recently published an account of their work in the June 3, 2007 issue of *Nature Materials*. The researchers have been able to determine the gene sequence for reflectin and use *Escherichia coli* bacteria to study how the protein grows under different temperatures and chemical conditions. The photonic properties of the variously synthesized reflectin molecules were then studied in detail. The key point is that reflectin proteins can self-assemble into a number of different configurations, observable on a nanometer scale (discussed further in Chap. 5), with a concomitant change in their reflective properties. The researchers have been able to produce solutions of reflectin that can be spread upon silicon wafer substrates and they ultimately hope to manufacture reflectin-based metamaterial substrates. Such substrates, of course, could have useful application in invisibility cloaking and active camouflage devices.

Eaton's Singular Cat's Eye Lens

Is there such an object as a perfect reflector? In principle the answer is yes, but in practice it is not easy to assemble. The properties of the cat's eye reflector were discussed earlier in this chapter, and its function is to return light rays back in the direction from which they came. The cat's eyes found in retroreflectors and road markings are made of material with a constant refractive index and are therefore only an approximation to the cat's eye ideal.

In order to make a perfect cat's eye reflector the refractive index inside of the sphere must vary continuously – it is indeed a GRIN lens, as first dreamed of by Robert Hooke. American mathematician J. E. Eaton (while at the Naval Research Laboratory in Washington D. C.) described the refractive index variation required to construct a perfect cat's eye reflector in the early 1950s – if the reflector is a sphere of radius R , then the refractive index must vary as $n(r) = \sqrt{2(R/r) - 1}$, where r is the distance from the lens center.

While the formula for the variation of the refractive index looks simple enough it actually contains a fatal, in practical construction terms, that is, requirement. It turns out that near the center of the sphere where r is getting closer and closer to zero so the refractive index becomes larger and larger and it is technically infinitely large at the very center. The Eaton, perfect cat's eye, GRIN lens (Fig. 3.29) cannot be physically constructed because of this central singularity (where theoretically the speed of light must go to zero). One can fabricate an approximate Eaton lens, however, by breaking the refractive index variation into a so-called step function (this, recall, was Robert Hooke's idea). That is, the lens is constructed layer by layer as a

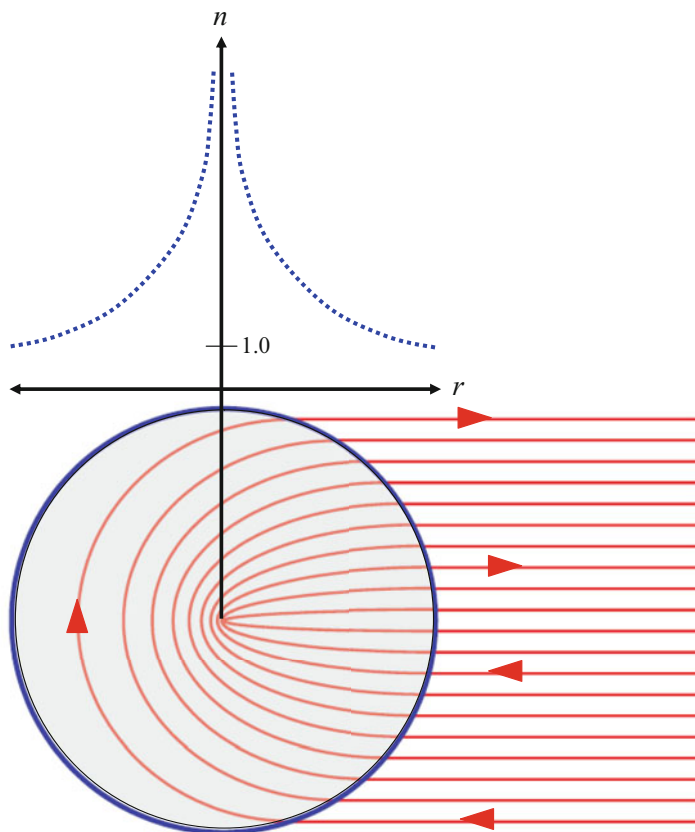


Fig. 3.29 Ray paths for light passing through an Eaton lens. Since the refractive index profile is radially symmetric the retroreflection will work from all incident directions. The schematic variation of the refractive index $n(r)$ is shown by the dotted line

series of smoothly joined concentric spheres, each of finite thickness, and each with a different but constant value of refractive index. This fabrication technique still poses, however, a technical challenge, and no true Eaton lens, working at optical wavelengths, has ever been produced.

An approach to the fabrication of a perfect Eaton lens, but doing away with the concentric glass shells, has recently been described by Tomas Tyc and Ulf Leonhardt (University of St. Andrews, Scotland). The new device is made of a specially constructed metamaterial (to be described more fully in Chap. 5), and while the problem of dealing with the central singularity still exists, there are ways of minimizing its effect. The key idea is to use specially produced metamaterials that guide the light rays around the lens as if there really was a singularity at the center. The trick, in practical terms, is to work out exactly how the metamaterial must be constructed to mimic the effects of the singularity, without there actually being a physical singularity at the center.



Fig. 3.30 A two-dimensional Eaton lens for microwave radiation. The device is composed of a metamaterial made of copper split-ring resonators etched onto circular mica circuit board substrates. The entire device is 80 mm in diameter (Image from arXiv:0906.5214v1)

While Tyc and Leonhardt have shown that a perfect Eaton lens operating at visual wavelengths could be constructed in principle, no such device has yet been manufactured. This being said, a microwave Eaton lens has been tested by the St. Andrews University group with collaborators from Singapore and the Czech Republic. The two-dimensional microwave Eaton lens (Fig. 3.30) is described in the journal *Nature Materials* for June of 2009, and it is a metamaterial made up from ten circular substrates composed of printed circuit board material. The rings are 10-mm high and set concentrically; each ring is 4 mm larger in radius than the next one interior to it. The entire device is 80 mm in diameter. Etched onto the printed circuit board substrates are a series of copper split-ring resonators (we shall describe the detailed workings of these later in Chap. 5). The purpose of these resonators, and an infilling of dielectric powder between the substrates, is to provide the device with the required refractive index profile.

Maxwell's Super Fisheye Lens

A variant of the Eaton lens is that of Maxwell's fisheye. While there is nothing actually piscatorial about this lens, it is a lens with many remarkable features—indeed, it is a lens that in principle can pick up and transport image detail that is invisible to other lenses.

James Clerk Maxwell introduced his fisheye lens in terms of a purely theoretical problem. Specifically, it was a problem he anonymously set out to challenge the esteemed readers of the May 1853 issue of the *Cambridge and Dublin Mathematical Journal* with. What Maxwell asked is, can a spherical lens be fabricated such that every light ray emanating from a point P on its surface converges to another point P' , also on its surface, with P' being the conjugate point of P . This all sounds very

technical (and it is), but the upshot of the question is that, yes, such a GRIN lens can be envisioned, and the refractive index within its interior must vary with distance r from the center, so that $n(r) = 2 / (1 + (r/R)^2)$, where R is the radius of the sphere. No true, fully functional Maxwellian fisheye lens has ever been crafted⁹ (or grown for that matter), but the potential properties of such lenses continue to fascinate physicists and mathematicians alike.

One of the more remarkable features of Maxwell's fisheye lens has only recently been discovered by Ulf Leonhardt (St. Andrews University). Writing in the September 2009 issue of the *New Journal of Physics*, Leonhardt showed that Maxwell's fisheye is in fact a superlens. We shall have more to say about superlenses (and hear much more from Ulf Leonhardt) in Chap. 5, but for the moment the point is that Leonhardt has shown that Maxwell's fisheye lens can, at least in principle, resolve image detail with unlimited precision. This fine resolution feature is something that cannot be achieved with standard optical lenses. Indeed, the so-called diffraction limit restricts the resolution of ordinary lenses to be no better than the wavelength of the light illuminating the object.

The Luneburg Cloaking Array

Another variant on the Eaton and Maxwell fisheye theme is that of the Luneburg lens. Again, it is a GRIN lens with special light-manipulating properties, and recent research suggests that arrays of such lenses might make an effective invisibility cloak.

The Luneburg lens was introduced during a graduate course in optics at Brown University in 1944 by the German-born mathematician Rudolf Luneburg. Like the Eaton and Maxwell fisheye lenses, the Luneburg lens requires the refractive index to vary with interior radius r . In the Luneburg case the variation is $n(r) = \sqrt{2 - (r/R)^2}$, and importantly, in contrast to the Eaton lens, this equation has no singularity at the center – the refractive index simply varies between the central value of $n=1.414$, where $r=0$, and a surface value of $n=1.000$ when $r=R$. While no substantive Luneburg lens has been crafted to operate at optical wavelengths¹⁰ they have been successfully made to function at longer radio, microwave and radar wavelengths. A ray diagram revealing how a Luneburg lens works is shown in Fig. 3.31, and it can be seen that all the light rays incident upon one hemisphere are brought to a single focal point on the opposite hemisphere. Indeed, it is this very property that makes such lenses (usually constructed from polystyrene plastics) useful as universal receiving antennas and radar navigation reflectors.

⁹This being said, as we shall see in Chap. 5, within the past several years two-dimensional analogs of Maxwell's fisheye lens have been fabricated.

¹⁰As with the Maxwell fish-eye lens, the first two-dimensional analogs of a Luneburg lens have only recently been produced. The first such devices are actually miniscule in size and are fabricated using techniques borrowed from the electronic microchip industry.

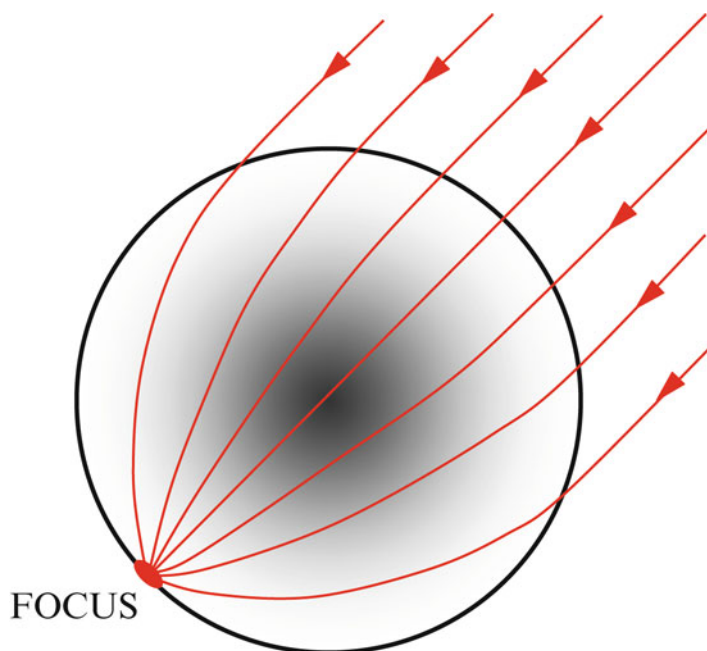


Fig. 3.31 Ray diagram for a Luneburg lens. The shading indicates that the refractive index increases towards the center of the sphere

Using Fig. 3.31 as a guide, imagine what will happen if two Luneburg lenses are brought together such that their surfaces just touch at one point. The effect would be that parallel rays entering the first (receiving) lens will be directed towards its focal point, and from there transmitted to the focal point of the second (radiating) lens, from which they will exit and continue on their way as if neither of the lenses had been encountered.

Expanding this idea to form a two-dimensional array of Luneburg lenses is one way of making an optical cavity – or an invisible cave, if you will. Indeed, research conducted by Niels Mortensen and colleagues at the Technical University of Denmark, and published in the March 4, 2009, issue of the *Journal of the European Optical Society*, describes how an array of Luneburg lenses might be used to make a microwave, or radar cloaking device. In this case, as illustrated in Fig. 3.32, the material to be cloaked is hidden in the pore space between the sphere surfaces. The size of the invisibility cavity, in which solid objects can be hidden from direct view, will vary in step with the diameter of the lenses used in the array. This size constraint might not be a problem for smaller objects, but it is clearly not a practical means of cloaking large structures such as buildings and/or aircraft.

While two osculating Luneburg spheres are in essence invisible (the light ray paths are the same before and after), it is interesting to ask if a single spherical lens can achieve the same effect. The answer to this question is, “Yes it can,” and the

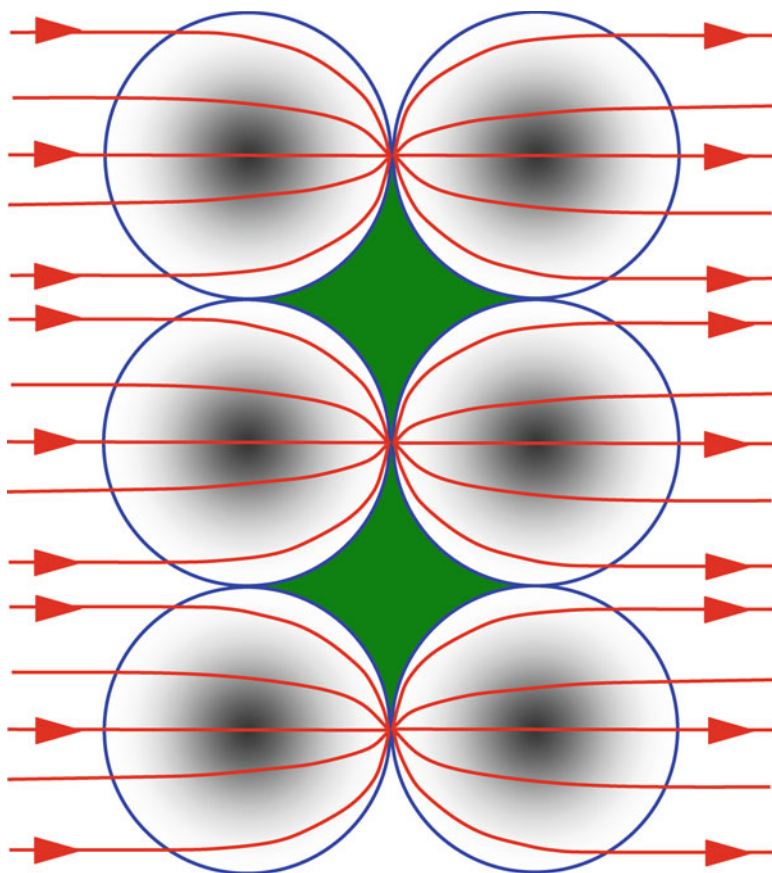


Fig. 3.32 A multiple Luneburg lens cloaking array. Any objects placed between the spheres (*dark region*) will be invisible at optical wavelengths to an outside observer

details on how have been worked out by Tomáš Tyc and Ulf Leonhardt. The method, however, is none too trivial and relies upon a complex variation in the internal refractive index. Indeed, the incoming light rays follow a path that loops around the center of the sphere (Fig. 3.33). While the sphere is technically invisible, it cannot be used to hide anything within its interior since the light rays must pass through it and not around it. The invisibility of the sphere could, in principle, be betrayed by timing how long it takes a light ray to travel between two points, A and B (see Fig. 3.33). The rays passing through the sphere, on path A' to B', will take longer than those running straight from A to B, since they have to loop around the sphere's center. In practical terms an invisible sphere could be used as a light-ray delay sequencer – to mix our metaphors, while time is money, a time delay sequencer is, apparently, invisible.

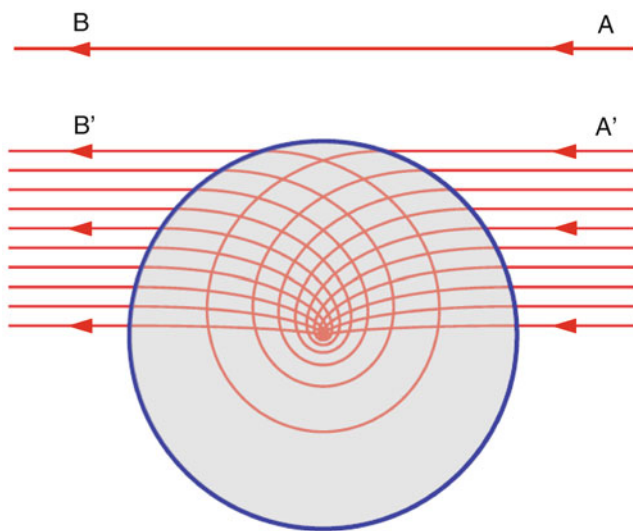


Fig. 3.33 Light rays passing through an invisible sphere. The light travel time from A to B will be shorter than that from A' to B' because of the shorter distance traveled

A Crystal Invisibility Cloak

Calcite is a clear, naturally occurring carbonate mineral, and it has long been known to have strange optical properties. Indeed, calcite is a so-called birefringent material. First studied by Danish scientist Rasmus Bartholin in the late 1660s, calcite is different from other optical media, such as glass, in that when light passes through its surface it is split into two rays, each of which is diffracted slightly differently. This effect comes about because of the arrangement of the atomic planes within the crystal and the fact that it has a different refractive index according to the polarization plane (recall Chap. 2) of the incoming light.

These special properties of calcite have recently been exploited to develop what are described as carpet cloaks by two independent research groups, one headed up by Xianzhang Chen (University of Birmingham, England) and the other by Baile Zhang (Singapore-MIT Alliance for Research Technology – SMART – Center, Singapore). The trick, it turns out, with these invisibility cloaks is to use a linearly polarized light source and two carefully crafted calcite prisms glued back to back.

The idea behind a carpet cloak is to make a surface having wrinkles and gaps in it (where objects might be hidden) look like a perfectly flat mirror to an external viewer. In essence the optical cloak smoothes out the bumps so that any light ray incident upon the cloak is reflected as if it encountered a flat mirror. The calcite prism cloak developed by the SMART group is described in the January 21, 2011, issue of *Physical Review letters*, and Fig. 3.34 shows it in operation, hiding a thin

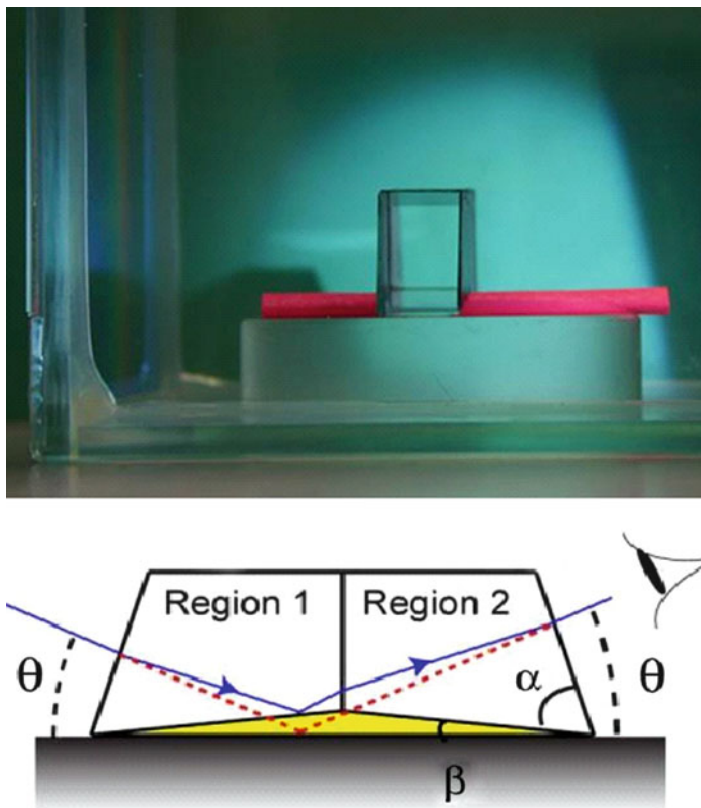


Fig. 3.34 (Top) A region of red card disappears from view when it passes underneath two calcite prisms. The entire prism is actually immersed in a light-oil bath since the eventual idea is to use the design as an underwater cloaking device where light is naturally polarized. (Bottom) Schematic light ray path through the calcite prism (Image courtesy of SMART)

piece of colored card. The calcite crystals are just 38 mm long, 10 mm wide and 14.5 mm in height, and they produce a triangle cavity of height 2 mm. Anything within the triangular cavity will be invisible to outside view, and the presence of the cavity will be hidden since any light passing through the cloak will behave as if it was reflected off a perfectly flat surface.

The calcite cloak developed by the University of Birmingham group is described in an article uploaded to the arXiv website on December 13, 2010. Their triangular prism arrangement, prosaically called a “macroscopic volumetric cloaking device,” is 20 mm high and can conceal a triangular cavity that is 1.2 mm high.

Both the SMART and the Birmingham cloaks open the way for larger, more practical optical cloaks, but they will never hide large structures such as buildings and vehicles – the crystal prism structures would be just too big and too obvious to prying eyes. Not only this, the cloaks only work according at one angle of polarized light, and they are also only two-dimensional. In the latter sense, if you view the

prisms from the side then there is no cloaking effect at all – you will see both the prisms and the cavity. But it is a start, and calcite is a relatively easily formed and engineered substrate. In Chap. 5 we will take a look at several more carpet cloak designs, but these will be made of specially fabricated metamaterials rather than natural mineral prisms.

The Cone of Silence

Since sound is a wave-like phenomenon it has many characteristics that parallel light. It is natural to ask, therefore, if the auditory equivalent of an invisibility cloak can be made – literally a cone of silence. Indeed, there are many real world applications where acoustic cloaking devices would be very useful: to hide a submarine, for example, from sonar detection (sonar – SOund Navigation And Ranging – being just short pulses of sound moving through water) and in the design of auditoriums and sound studios. In the latter case, for example, acoustic metamaterials could be used to either absorb sound, from say a badly reflecting corner, or to directly focus sound waves towards a location where it is needed.

Several research teams are working on the development of acoustic cloaking devices. To date only a few prototypes and concept models have been constructed, but the theoretical details have been worked out in detail. One team of researchers, Steven Cummer and David Schurig at Duke University (Fig. 3.35), has performed a series of complex computer simulations that follow the interaction and reflection of

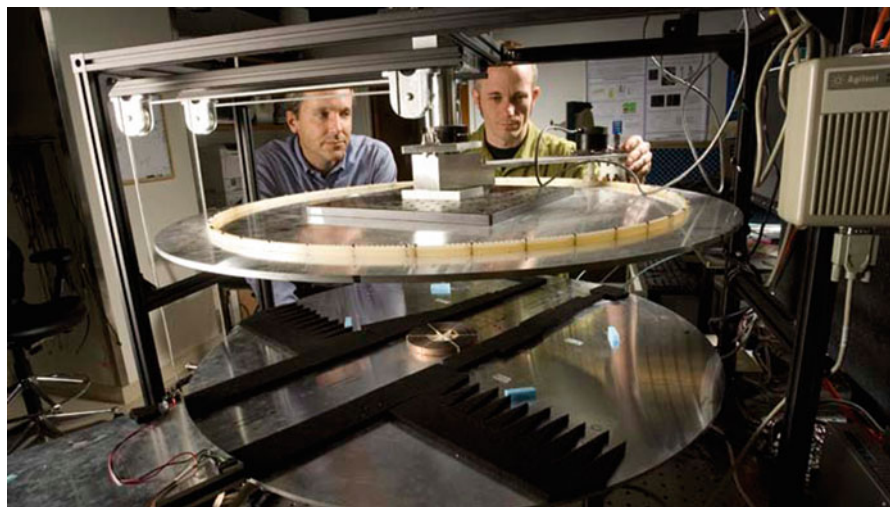


Fig. 3.35 Steve Cummer (*left*) and David Schurig (*right*) making adjustments to the acoustic cloak testbed. A microwave invisibility cloak (to be described in Chap. 5 – but recall Fig. 3) can be seen on the lower disk of the testbed (Image courtesy of Duke University)

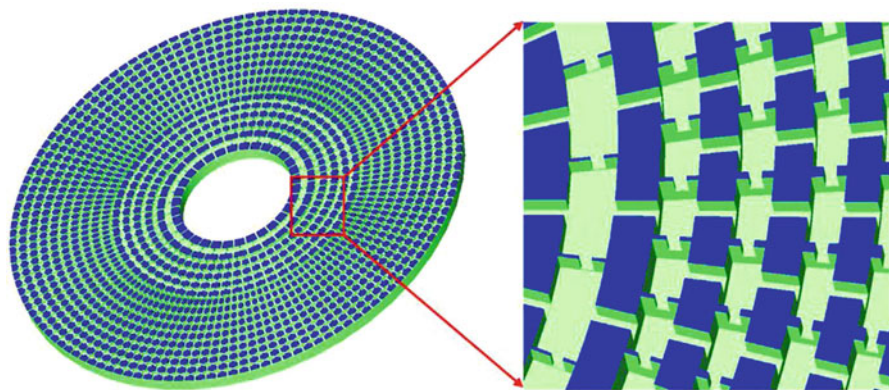


Fig. 3.36 An acoustic cloaking device. The concentric cylinders and connecting channels (*inset*) prevent an object located in the central cavity from being detected by sonar (Image courtesy of Nicholas Fang, University of Illinois at Urbana-Champaign)

sound waves within materials of different density and elasticity. In this manner they have developed a description of how a spherical acoustic cavity might be made – at least in principle. It will be no simple task to construct a sound invisibility zone, and no naturally existing materials have all the right properties. The acoustic shell will have to be carefully manufactured, layer by layer, with very specific, built-in mass density and bulk modulus characteristics varying with distance from the region to be cloaked. The point here is that by adjusting the cloak's material qualities the sound speed through it can be controlled and the sound waves smoothly wrapped around the acoustic cavity. One way to manufacture an acoustic cloak might be to use a composite structure made of sonic crystals – materials with periodic variations in their density (to be describe shortly).

The first practical acoustic cloaking device was recently developed by Nicholas Fang and co-workers at the University of Illinois at Urbana-Champaign. The demonstration device (Fig. 3.36) is composed of 15 concentric cylinders with an outside diameter of about 10 cm. The cylinders contain specially engineered cavities and channels that direct sound waves around a central cavity some 25 mm across. The prototype device was designed to work underwater at ultrasonic frequencies and is therefore a sonar invisibility cloak. While such technology can in principle hide objects from underwater detection they are bulky, and the object-to-cloak size ratio currently prohibits them from being effective cloaking devices for, say, a moving submarine.

Sonic Crystals

Sounding like something straight out of a science fiction fantasy story, sonic crystals are very much real entities, and they are an example of a metamaterial specifically manufactured to manipulate sound waves. The first such crystal was described

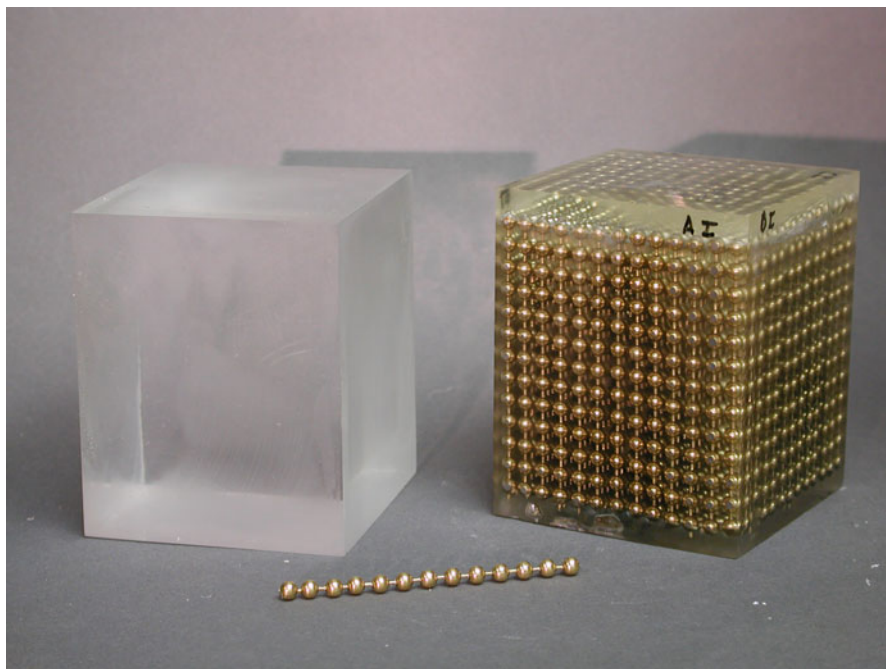


Fig. 3.37 A sonic crystal composed of small copper spheres embedded within an epoxy resin matrix (Image courtesy of the U.S. Navy)

by Zengyou Liu (Hong Kong University of Science and Technology) and co-workers in a research paper published in the journal *Science* in 2000. As the name suggests, sonic crystals have a lattice structure, and in the 2 cm slab of material described by Liu et al., the lattice was made of strings of small lead spheres that had been coated with silicone rubber. The chains of spheres were then embedded in an epoxy resin matrix (Fig. 3.37). It is the periodic variation in the density of the material, due to the spheres, that robs energy, at very specific wavelengths, from the sound waves as they propagate across the crystal. The spacing and size of the spheres can be adjusted to fine-tune the specific wavelengths that a given sonic crystal absorbs.

While a given sonic crystal can inhibit the passage of sound waves at specific wavelengths, a more complex arrangement of different kinds of sonic crystals will be required to produce a full acoustic cloaking device. Spanish researcher Jose Sanchez-Dehesa (Polytechnic University of Valencia) and co-workers have, however, taken on the challenge of building a practical acoustic cloak. Writing on their work in the *New Journal of Physics* in 2008, Daniel Torrent and Sanchez-Dehesa describe in detail a two-dimensional cylindrical cloaking system made up from multiple rings of sonic crystals. Indeed, building upon the computations by Cummer and Schurig (Duke University) the Valencia group argue that a practical acoustic cloaking device can be made from just two kinds of equal-thickness sonic crystals arranged in concentric shells (Fig. 3.38).

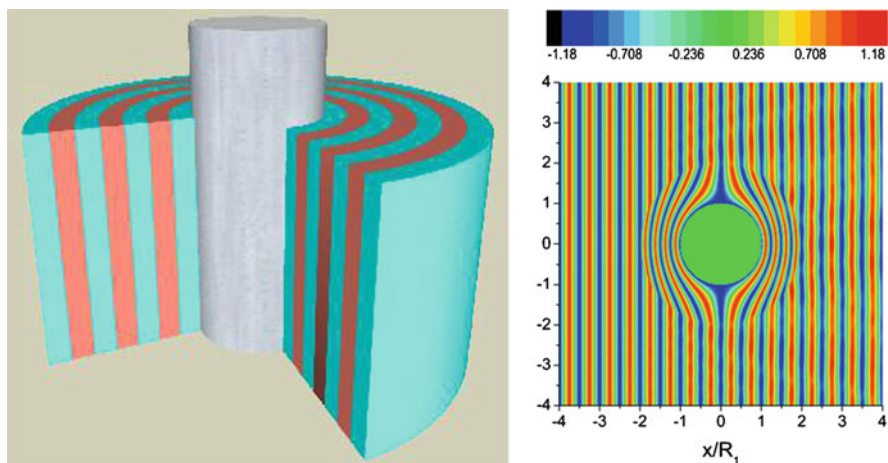


Fig. 3.38 (Left) An acoustic cloaking cylinder made up from alternating layers of two distinct types of sonic crystals. (Right) Pressure map of a sound wave striking the acoustic shell and wrapping, as required, smoothly around it. In this example, the radius of the outer shell is twice that of the inner shell (Image courtesy of the *New Journal of Physics*)

A full working model of the acoustic cloak designed by Torrent-Sanchez and Dehesa has yet to be realized, but as with optical invisibility cloaks, we are on the very threshold of seeing this new technology being utilized and put to practical use.

A Wave-Less Sea

Of all the many different types of waves that surround us the ones that look most, well ... wave-like, are those that ripple across the surface of a pond, or crest and then pound the shoreline by the sea. Indeed, the waters of the oceans heave and swirl with relentless tenacity, enacting a primordial battle between the elements, and eroding the coastlines of the world.

Water waves, just like sound and light waves, carry energy, and that energy can be dangerous to all manner of maritime objects, from coastal barricades, to boats on the sea, and off-shore oil-platforms. Emerging from the new research relating to invisibility cloaking, however, is a possible defense mechanism against rogue water waves and on-shore erosion. Just as with acoustic cloaking, the key to the process is to systematically change the speed of the waves so that they move around the cloaked device, leaving a central sweet spot where there is either a millpond surface, or at least an extended region over which the wave energy is significantly reduced.

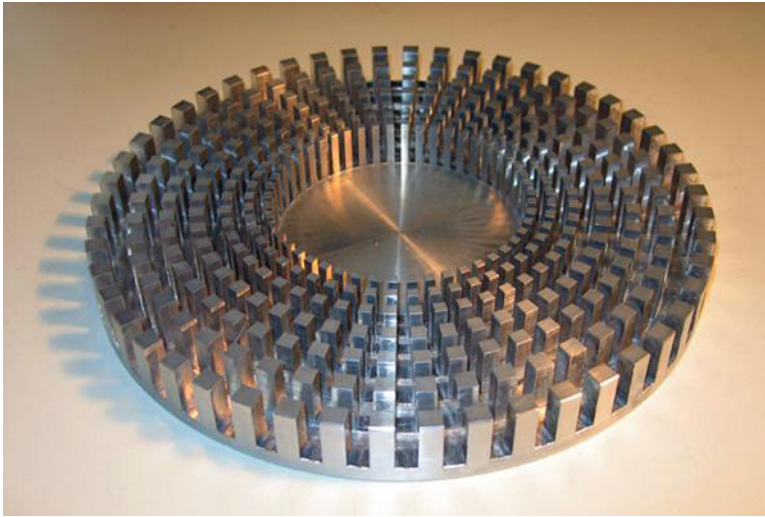


Fig. 3.39 Water wave cloaking device. This concept proving device is 20 cm in diameter evenly machined into seven concentric rings and 100 rigid sectors. The inner, 8.2 cm diameter region is cloaked from outside surface waves

The detailed design of a water wave cloaking device was presented by Stefan Enoch (University of Aix-Marseilles, France) and Sebastien Guenneau (University of Liverpool) in September 2008 in the journal *Physics Review Letters*. The concept demonstration device is composed of a flat circular plate with a series of concentric protruding sectors (Fig. 3.39). The key idea is to guide the water waves around and between the sectors such that their velocity is greater in the direction around the circumference of rings than it is in the radial direction between them. Figure 3.40 illustrates this effect, and it can be seen that the surface wave crests (dark bands) emanating from a generating source (upper right in the diagram) move around the cloak (lower left), leaving the central region (white circle) undisturbed.

The efficacy of the water cloak has been established by detailed numerical simulations (as in Fig. 3.40 above) and also in laboratory experiments. To date no life-sized cloak has been constructed, but in principle, Enock and Guenneau comment, their design could be scaled-up to protect deep-sea oil platforms from damage due to rouge waves, or protect lowland coastal regions from the deadly effects of a tsunami (at least localized regions anyway). Semi-circular sections of the cloak could also be constructed to help protect small islands, beaches and shorelines against wave erosion. Almost as an afterthought, in the final paragraph of their *Physical Review Letters* article, Enock and Guenneau also note that their water wave cloaking design could be used as an acoustic cloak, capable of constraining so-called shear waves that form along cracks in flexible solids. In this latter application the cloak could act to protect a structure, such as a building, against integrity damage from hurricanes, tornadoes and even earthquakes.

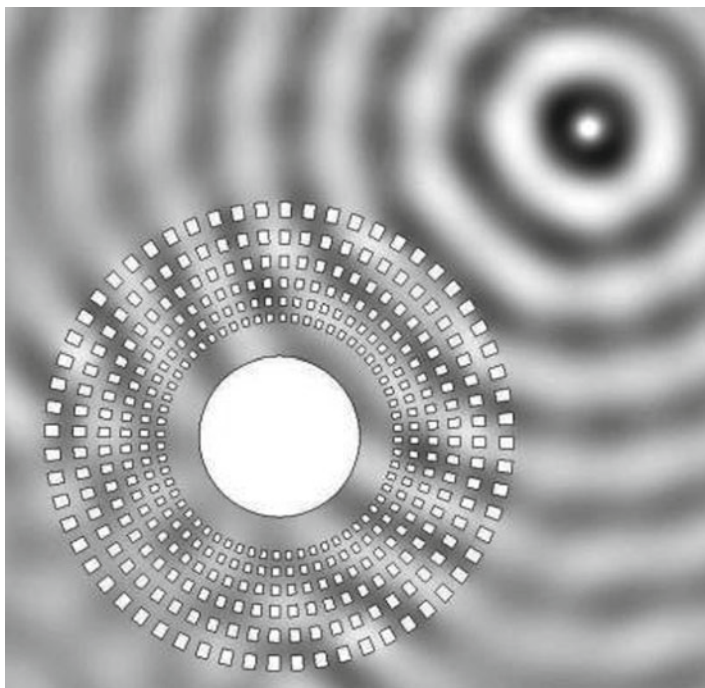


Fig. 3.40 Simulations showing the circular surface wave crests moving around the invisibility cloak, thereby producing an inner protected wave-free region (*white disk*)

And Earth Stood Still

Building upon the idea of water wave cloaks, it seems natural to ask if it might be possible to protect buildings against earthquakes? The answer is yes, in principle, and again it was Stefan Enock and Sebastien Guenneau who paved the theoretical groundwork.

Earthquakes produce two kinds of waves that can cause structural damage. One is a body wave that passes directly through the surface rocks while the other is a surface wave that ripples along the ground. It is these latter surface waves that cause the most structural damage during an earthquake, and if a cloaking device against them could be constructed then much damage and death might be avoided. Writing in the journal *Physics Review Letters* for July 2009, Mohamed Farhat (University of Aix-Marseilles, France) in collaboration with Enock and Guenneau presented a theoretical earthquake cloaking metamaterial device composed of 10 concentric rings made of various kinds of plastic. Each layer in the cloaking collar has a different stiffness and elasticity, selected so that when working in concert any impinging surface wave from the earthquake is directed around the cloak. The cloaking device doesn't destroy or damp-out the earthquake wave, so it will carry on and possible

cause damage elsewhere, but it does protect the buildings located within its interior from damage.

The earthquake cloaking device should ideally be built into the foundations of a building, and Guenneau comments that to protect a 10-m sized building, the various rings would need to be about 10 cm thick and have diameters between 1 and 10 m. Smaller versions of the cloak could also be deployed to control vibrations in manufacturing machinery, aircraft and automobiles. The next step in the study of earthquake cloaking is to move from the theoretical domain to physical buildings – a process that has yet to be realized.

In this chapter some of the ways in which light (sound, water and seismic waves) can be tricked, manipulated and distorted have been described. Next, however, we need to catch up with the unfinished business articulated at the end of Chap. 2. Indeed, we have seen how light propagates, but we have not as yet explained what light actually is.

Chapter 4

Maxwell's Marvelous Waves

Chapter 2 closed with the question, “What is light?.” This seemingly harmless-sounding question will be answered in this chapter, but it is an answer that took scientists many centuries to fully articulate. The story behind the answer, however, begins by looking at what at first glance seem to be two totally unrelated phenomena – the phenomena of magnetism and electricity. For indeed, to understand light we need first to understand the riddle behind these two optically invisible entities.

Magnetic Attraction

Dr. William Gilbert saw the first copies of his *De Magnete* – literally, “The Magnet” – appear in early 1600. It was a book full of mysticism, exacting experimentation and vented bile. Gilbert was specifically interested in magnets and especially the idea that Earth itself was a giant magnet. His experiments with terrella (Fig. 4.1), which are small spherical lodestones (naturally magnetic rock containing magnetite), convinced him that not only was Earth a magnet but that it was in a very real sense alive, with the traveler-guiding magnetic field being a tangible expression of its benevolent soul. It was a radically new idea.

For all his odd ideas about the origin of Earth's magnetic field, Gilbert did have a point. There is indeed something primordial and spooky – and fascinating – about magnetism. Pick up any two magnets (your fridge at home is probably covered with them), hold them in your hand, and look at them. They are apparently inert, cold, dark pieces of metal and yet bring them together and they either attract or repel each other – seemingly out of nowhere an invisible force materializes between the advancing magnets, drawing them together with a snap, or pushing them rudely apart.

There is absolutely nothing in the external appearance of the magnets that would have hinted at, or even predicted, that they would behave as they did. Like magic their close proximity seemingly brought into existence a new, invisible, dynamic, otherworld entity. Magnetism is an ancient mystery, and it is not known who first

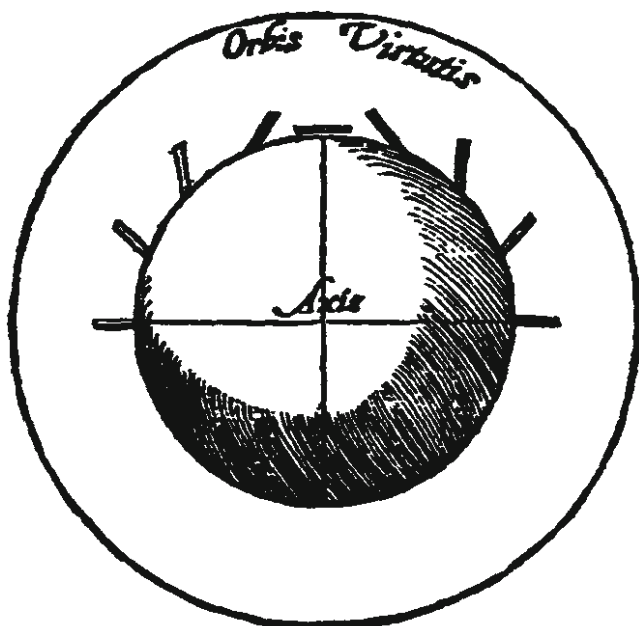


Fig. 4.1 Diagram from Gilbert's *De Magnete*. Results from experiments conducted with a terrella. The bars shown on the terrella's surface indicate the direction of magnetic inclination. The outer circle reveals the *orbis virtutis* (the sphere of virtue) that delineates the region of the terrella's magnetic influence

noticed its effects. Neither is it known who first observed that the lodestone, a naturally magnetic material (Fig. 4.2), will turn if freely suspended and point northward – a natural aid to navigation and the savior of many a lost explorer.

While William Gilbert certainly brought the experimental study of magnetism to a new level of sophistication, and although the magnetic nature of Earth and other magnets could be mapped out with needles and thin wire threads, it was still far from clear what magnetism actually was. Indeed, after the publication of *De Magnete*, a further 200 years would pass before the first inklings of a correct theory for magnetism would appear.

Laying on the Leyden Jar

Static electricity is a strange, invisible force that can be imparted to otherwise benign objects. Not surprisingly, it caught the eye of William Gilbert, and in addition to describing magnetic phenomena in his 1600 *De Magnete*, he also studied the amber effect, or as we would call it static electricity.

Although his studies on static electricity occupied just a minor fraction of his work, Gilbert carefully explained, correctly as it turns out, that the attractive force



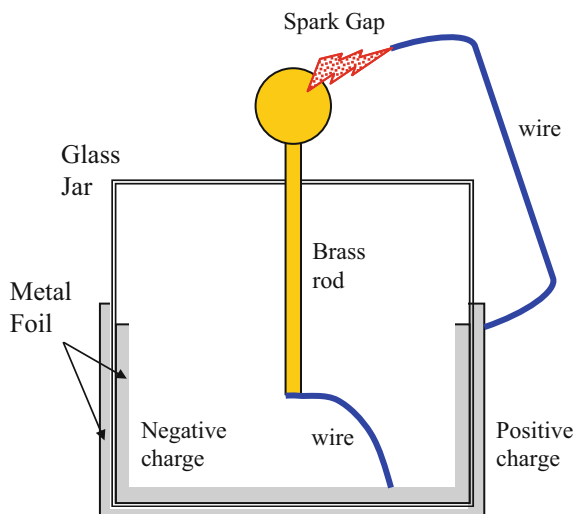
Fig. 4.2 A lodestone sandwiched, as if it might inexplicably fly away, between two iron plates. A strong magnetic field is established between the two vertical vanes (Image courtesy of Westminster College, Pennsylvania: <http://physics.kenyon.edu/EarlyApparatus/index.html>)

due to static electricity was very different in form to that of magnetism. The attractive effect that Gilbert experimented upon was produced by rubbing an amber rod with a fur cloth. Indeed, the word *electron* (leading to the word *electricity*), derived from the Greek word *electron* (meaning amber) was coined by Gilbert in his 1600 text. Amber, a hard, translucent, yellow-colored fossilized resin, was well known for its ability to be transformed by vigorously rubbing its surface. The rubbing mysteriously infused the amber with an invisible attracting force. What was an otherwise inert, non-interactive body could be magically transformed into a paper- or wood pith-lifting corpus. The same effect has often enlivened many a child's birthday party when a balloon that has been rubbed against a woolen jersey can be made to stick to a wall or the ceiling – simple magic indeed.

In modern terms the phenomenon we are dealing with is that of triboelectricity. This is the process by which an electric charge is produced on the surface of a dielectric, via mechanical separation, after rubbing it with a cloth. A dielectric is any material that is normally an insulator – which is a material that does not allow electricity to flow through it. A conductor, in contrast to an insulator, is a material that readily allows electricity to flow along its length.

Experiments performed by American polymath Benjamin Franklin in the early 1750s revealed that two kinds of electrical charge could be produced depending upon the substance being rubbed and the type of cloth being used: a positive charge was produced by rubbing glass with a silk cloth, while a negative charge could be

Fig. 4.3 Charging a Leyden jar. The glass container acts as the insulating region between the negatively charged inner foil and the positively charged outer foil. The charge is released by making a connection between the outer foil and the brass rod connecting it to the inner foil



produced by rubbing sealing wax with a fur cloth. It was also apparent that, just as with magnets, objects supporting like charges repelled one another, while objects supporting opposite charges attracted one another. The struggle to determine the exact nature of electricity, however, occupied many centuries of further study.

In spite of having a long history, the amber effect was generally considered to be little more than an amusing play thing – it was an oddity, and it was temporary, the attractive effect being lost after just a few minutes. This all began to change in 1745 when in the Dutch city of Leyden a means of storing and building up large electrical charges was developed. The device was invented by Pieter van Musschenbroek and is commonly known as the Leyden jar.

The Leyden jar is essentially a glass container with its inner and outer surfaces covered with metal foil. The inner foil surface is attached to a small wire or chain leading to a charging sphere. The jar is primed by allowing a charge to build up on the inner foil (Fig. 4.3). As the charge builds up on the inner surface of the jar, say as a negative charge, the outer metal foil, which is grounded during the charging process, will build up a positive charge. This charge separation comes about because like charges repel each other, while opposite charges attract. In effect, the negative charges in the outer foil are driven down the ground wire and away from the increasing negative charge on the inner foil.

With the separate charges established on the inner and outer foils, the grounding wire connected to the outer foil is removed. The Leyden jar is now primed and ready for discharge – and it can deliver quite a kick. Musschenbroeck wrote to French physicist Abbe Nollet after receiving an accidental discharge that, “I thought it was all up for me... I would not take a second shock for the kingdom of France.” The jar is discharged by simply bridging the gap that connects the inner and outer foils. The name Leyden jar was in fact first used by Nollet, and it was also Nollet who performed what is perhaps the ultimate in public scientific demonstrations when in

1746 he discharged a Leyden jar spark through 180 of King Louis XV Versailles guard – making them, to the king’s great delight, leap into the air simultaneously. It was a shocking display of electricity’s raw power, but again it was nothing more than a public plaything and a means of childish amusement.

The properties of the Leyden jar were first clearly articulated by Benjamin Franklin in his *Experiments and Observations on Electricity*, published in 1751. Franklin showed that the charge was not only distributed over the inner and outer foils, but that it was also proportional to the surface area of the foil. In modern terms the Leyden jar acts like a capacitor, storing and releasing charge as required by the electrical circuit. Franklin, of course, is well known for his studies on the new science of electricity, and it was in 1752 that he performed his famous kite flying experiment, supposedly on a dark and stormy night, showing that lightning was an electrical phenomenon. It was from these experiments that he went on to investigate the so-called action of points and develop the very useful lightning conductor – a simple device that has saved many a building and countless human lives (Fig. 4.4). Indeed, Franklin was one of the new breed of natural philosophers grappling with the problems of a mercurial and dynamic science – it was a changing of the guard, and the investigation of electricity and magnetism was all the rage.

Michael Faraday

Faraday was very much the man of the moment. He espoused all the key qualities expected of a Victorian scientific superstar – he was deeply religious, raised from humble origins, self-taught, hard working, and a charismatic speaker. Faraday’s big break into the world of science, however, came when he was just 21 years old when he landed a job at the Royal Institution in London.

At the tender age of 34, in 1825, Faraday became director of the Royal Institution’s laboratory. This job required him to deliver the highly popular Friday evening lectures – well attended talks on topics relating to all areas of science and engineering. These talks were in many ways the Victorian equivalent of today’s reality TV programs. They attracted large public audiences and provided the attendees with a series of experiments and practical demonstrations for their greater edification and enjoyment. Savvy practitioners of their trade, the lecturers knew that experiments involving sparks and electrical phenomenon were a surefire crowd pleaser. Indeed, scientific instrument maker Benjamin Martin had long since observed in one of his many educational texts that electricity was, “an entertainment for angels, rather than men.”

Not only did Faraday excel in his lectures (Fig. 4.5) he also began to work on a number of pivotal experiments relating to magnetism and electricity. As for the sciences, Faraday first made his name in May of 1821. Having heard of a remarkable paper read by Danish physicist Hans Oersted at the Royal Institution in Copenhagen on October 1, 1820, that a compass needle could be deflected if passed close to a wire carrying an electric current, Faraday immediately determined to emulate and test the result. Within short order Faraday had not only confirmed Oersted’s



Fig. 4.4 A dramatic lightning flash illuminates the space shuttle launch pad. According to the Hartford Insurance Company, lightning is responsible for \$5 billion in insurance industry losses annually (Image courtesy of NASA)

result, he had also demonstrated the reverse effect, that a magnetic field could deflect an electric arc. These two experimental results demonstrated, for the first time, a definite symmetry between electricity and magnetism. Each phenomenon may well be explained according to different forces and emanations, but now, somehow, it appeared that they must have mutual sympathies.

With these ideas in place French physicist Andre Ampere became obsessed with the idea that magnetism was really just an expression of circulating electric currents and during the mid-1820s developed a detailed mathematical theory of electromagnetism. Today, Oersted's experimental result is generally known as Ampere's law

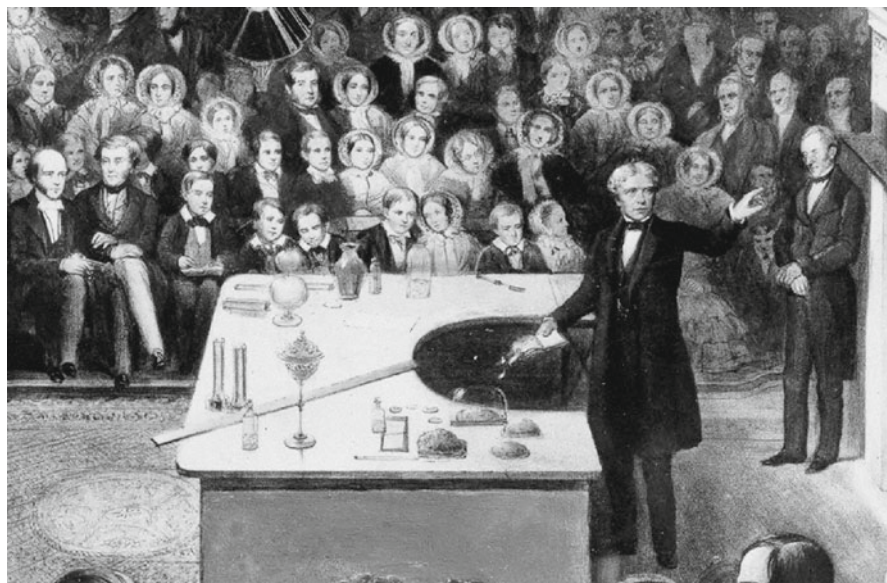


Fig. 4.5 Michael Faraday in full oratorical flow during the December 27, 1855, Christmas lecture at the Royal Institution in London. Note the large, enthralled, and all-ages audience, as well as the experimental apparatus arrayed on the table before him. Among the audience on this particular night were Prince Albert of Saxe-Coburg and Gotha, and the Prince of Wales (Engraving from the *Illustrated London News*)

and this, in its mathematical form, can be used to describe the magnetic field configuration associated with a specific closed loop of electrical current. In 1861 Maxwell added an additional term to Ampere's law to account for time varying electrical fields, but more on this later. Faraday's law is the complement of Ampere's law which, in its modern form, describes how a time varying magnetic field can induce an electrical current to flow in a closed electrical loop. In essence we have the result that moving electrical charges can generate a magnetic field (Ampere's law) and that moving magnetic fields can cause electrical currents to flow (Faraday's law).

Faraday brought light directly into the mix of electrical and magnetic field considerations when, in 1845, he discovered what is now appropriately known as the Faraday effect. Again, it was an experimental result that produced the idea – indeed, Faraday was no theoretician, and all his results were experiment based. Specifically, what Faraday found was that a magnetic field could alter the polarization properties of glass. The greater the magnetic field strength and the longer the path through the magnetic field, so the greater is the degree of rotation induced in the plane of linear polarization (recall the end of Chap. 2). It was a staggering and entirely new result, and it was the first hint that light itself must be some form of electromagnetic phenomenon.

Perhaps Faraday's greatest contribution to physics, after the discovery of his specific laws, that is, is his concept of electrical and magnetic fields – or field theory, as it would later be called. Indeed, during most of his life Faraday was

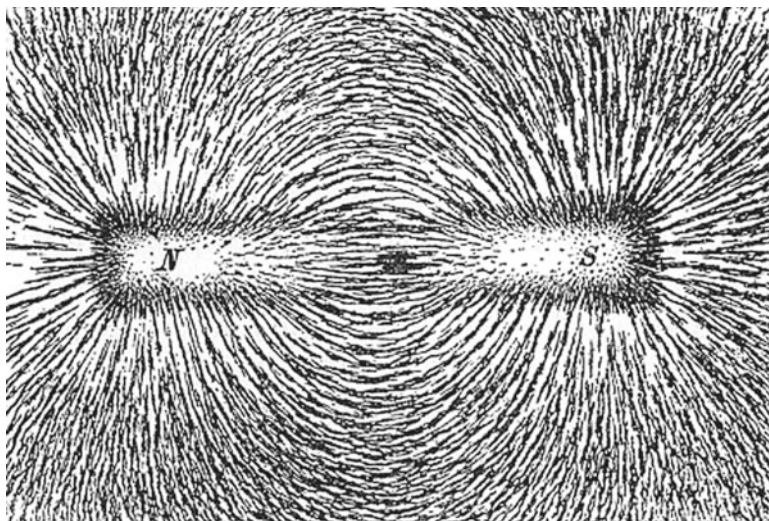


Fig. 4.6 A sprinkling of small iron filings reveal the otherwise invisible field lines emanating from a bar magnet (Image courtesy of Wikipedia commons)

chastised for sticking with his concept of fields of force, and the theoreticians of the day only tolerated his “childish” ideas because of his fame and clear brilliance as an experimentalist. The influence of a magnetic force, for example, Faraday argued, could be visualized according to a series of lines of force that filled the space around it. The direction of the lines could be displayed by an arrow, and the strength of the field could be judged according to how close the lines were packed together.

Figure 4.6 shows the experimental outcome of sprinkling iron filings onto a piece of paper placed over a standard bar magnet. This classic experiment is still performed in many school science classes, and it reveals the influence of the magnet upon its surrounding space. The magnetic field lines loop in various sized arcs trailing from the north to south pole locations of the magnet, and the strongest magnetic attraction (greatest field line packing) is at the north and south pole locations. An analog of Faraday’s magnetic field lines are the contour lines on a topographic map. Each contour on the map links together regions of the same height, and the closer the contour lines are to one another the steeper the rise or fall of the topography.

The field theory visualized by Faraday introduces a subtle shift in the apparent manner by which objects, such as magnets and electrical charges, interact with each other. The force felt between two magnets is not acting on their physical bodies, but rather it is an interaction between the field lines themselves. Faraday’s notion of magnetic and electrical attraction only became popular in the latter part of the nineteenth century when James Clerk Maxwell developed the appropriate mathematical tools for describing the field equations.

Faraday's Lightning Cloak

If Benjamin Franklin took some of the danger out of lightning strikes by advocating the installation of conducting strips of metal down the sides of buildings – to channel the electrical charge to the ground (Fig. 4.4) – then Faraday tamed lightning altogether.¹ Perhaps tamed is too strong a word, since the lightning still happens, but what Faraday essentially demonstrated was that lightning protection could be guaranteed to any observer located on the inside of a conducting metal cage.

Faraday's cage is in effect a lightning invisibility cloak, and it operates by quenching the buildup of any internal electric field – the idea being that if there is no electric field inside of the cloak then lightning cannot discharge there. The active process that makes the cloak work is charge redistribution within the cage's conducting substrate. Indeed, in the presence of an external electric field the freely moving electrons within the conductor will rearrange themselves (according to the like charges repel rule described earlier) so as to cancel out the development of a field inside of the cage.

The Faraday cage also works in reverse. In this situation the electrons within the conductor migrate to its inner surface so as to cancel out the electric field produced by any charges placed inside of the cage. In so canceling the inner field, however, a charge is built up on the cage's outer surface, but this excess can be siphoned off by grounding the cage. The exterior of the cage accordingly remains neutral in spite of the presence of an interior charge. It is by this cloaking effect that microwave ovens, for example, are rendered safe for kitchen use in that the microwave radiation used to heat the food does not escape from the oven's interior to literally cook the cook.

Gauss Lays Down the Law

Carl Friedrich Gauss (1777–1855) has been described as the prince of mathematicians. His interests were broad, varying from astronomy to number theory to electrostatics and optics, and in all these fields he left an indelible and magnificent legacy. He was a child prodigy, and his mathematical skills were exquisite.

In the early 1800s Gauss began to investigate the properties of magnets, but it was in 1835 that he derived the mathematical law that describes the electric field associated with a given arrangement of electrical charges. Specifically, Gauss's law describes the divergence of the electric field at each point in the field. We can think of the divergence as being a measure of the electrical field magnitude at every given point in space – that is, it describes how much of the electrical field enters a small region, about every point in space, and how much comes out – the outcome relating

¹ There is a certain amount of continued historical inaccuracy here in the sense that while Faraday certainly made experiments concerning charge distribution, the so-called Faraday's cage was initially described by Franklin.

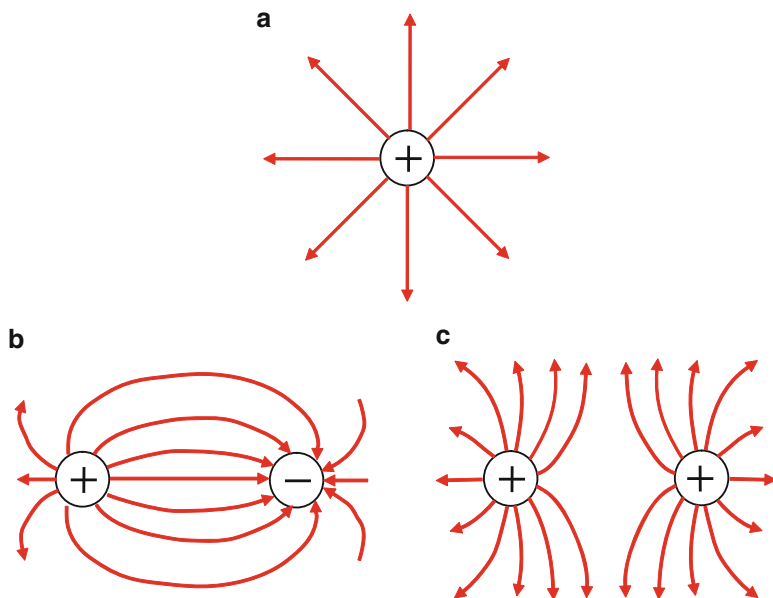


Fig. 4.7 Electric field line configurations associated with (a) a single charge, (b) two equal but opposite charges, and (c) two equal and like charges. The field line configurations can be mathematically deduced from Gauss's divergence law

to the charge distribution and what kinds of charges, positive or negative, are involved. As Faraday would say, it is a way of mapping out the electrical field lines associated with a specific charge distribution. Figure 4.7 shows the electric field line configurations produced by three simple charge distributions.

Although Gauss's law describes the electrical field line configuration associated with a given charge distribution, an equivalent law can be written down for magnetic fields. This latter law, however, is really just a statement concerning the non-existence of magnetic charges. Magnets, remember, are always found to have north and south poles, and there is apparently no magnetic equivalent in nature to isolated electrical charges. Well, that being said, the single pole magnet is in reality a beast of mythical proportions. Physicists genuinely believe that such entities really do exist, but to date no clear experimental detection of a monopole has been made. There are famous claims for detections, but no agreement upon the results – such is the way that science works.

A Beginning with String

James Clerk Maxwell (Fig. 4.8) is a true giant among the greatest scientists of history. Outside of physics, however, he is barely known and recently ranked only 91st in a British Broadcasting Company (BBC) public opinion poll for “the most famous Britons in history” – Maxwell was in fact sandwiched between King Henry II and



Fig. 4.8 James Clerk Maxwell with wife Katherine and faithful canine companion Tobi as photographed in 1869. Apparently Maxwell had several dogs throughout his life, all of which were conveniently called Tobi

J. R. R. Tolkien. World War II prime minister Winston Churchill took first place and Isaac Newton placed a respectable sixth. This admittedly none-to-scientific poll result is a remarkably sad placing given Maxwell's incredible intellect, and it runs counter to the claims made by many of his contemporaries that he was an equal to the great Isaac Newton. Perhaps saving some face, however, Maxwell ranked third behind Albert Einstein and Isaac Newton in a BBC millennium poll of 100 top physicists.

Continuing the parallel with Isaac Newton, Maxwell also studied at Trinity College in Cambridge, and like Newton he was somewhat of a recluse and loved nothing more than to be simply left alone – with a good mathematics problem to solve, that is. His first scientific paper was read to the assembled Fellows of the Royal Society of Edinburgh on April 4, 1846. At this time Maxwell was too young

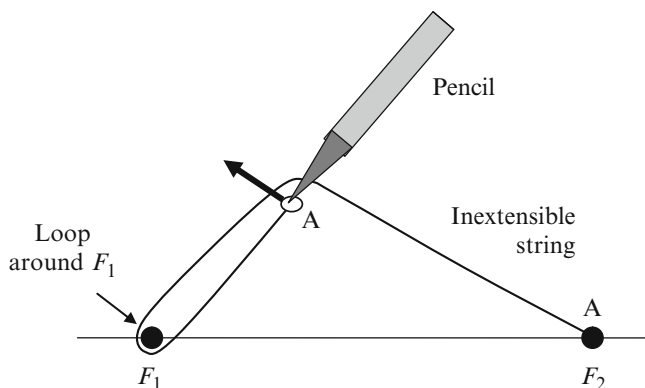


Fig. 4.9 Maxwell's method for constructing ovals. In this figure just one loop around F_1 is considered, but provided the string is long enough multiple loops around both F_1 and F_2 can be accommodated

to attend the society's meetings, being then only 15 years old and still at school answering to the nickname of Dafty. What a remarkable achievement. The paper itself concerned the drawing of ovals, and the method is so simple that once seen one's immediate response is, "Why didn't I think of that"? Well, of course, seeing the obvious before anyone else does is the signature of genius, and the young Maxwell had genius in abundance.

The method that Maxwell developed for constructing ovals of all shapes and sizes is illustrated in Fig. 4.9, and it builds upon and generalizes a much older method for drawing ellipses with a single length of string and two fixed pins. The two-pin method for drawing closed curves is, of course, the next step beyond the one-pin and string method that can be used to construct circles. What Maxwell found was that by looping the string around the second pin, rather than attaching it there, an oval, rather than an ellipse, could be drawn.

When he discovered his string and pins method for generating ovals the young Maxwell had no idea that it might be applied to optics,² and neither, presumably, did it occur to him that some 20 years later he would amaze the scientific world with the publication of his remarkable new theory on the constitution of light.

A Prelude to Electromagnetic Waves

Physicist Stephen Toulmin once observed that definitions are like belts – the smaller they are the more flexible they need to be. Fortunately, we do not need to worry about the full physical definitions employed in electromagnetism here; our belts

² It is in fact the analog computer solution to an otherwise tricky refraction problem that was first studied by Huygens in 1690.

will be short ones. The main reason for introducing a few formal definitions here is that they will enable us to make sense of Maxwell's great leap of intuition (to be described very shortly), and they will also enable us to understand what it is about metamaterials that allows them to act as invisibility cloaking devices and superlenses (topics to be covered in Chap. 5).

Our first definition relates to the term *dielectric*. A dielectric is essentially an insulator, or non-conductor, which means that an electric current won't flow through it. However, a dielectric is not totally blind to the presence of an electric field. By being exposed to an electric field the electrical charges within its constituent atoms will be very slightly displaced – the positively charged nucleus being displaced one way and the negative charge electrons the other. Through this action the dielectric is said to be polarized. This small shift and alignment of the charges within the dielectric creates its own small internal electric field that can then back-react with the external electric field, and so on.

In this fashion the dielectric modifies the external field $\mathbf{E}$ very slightly to produce what is called the displacement field $\mathbf{D}$, and the so-called electrical permittivity ϵ describes the relationship between the two, with $\mathbf{D} = \epsilon \mathbf{E}$. In a vacuum there is, of course, no material to polarize, and the electrical permittivity is equal to that of free-space and given by the symbol ϵ_0 . This same constant enters into what is known as Coulombs law, which describes the force acting between two charges a given distance apart. In System International (SI units) $\epsilon_0 = 8.85 \times 10^{-12} \text{ C}^2/\text{Nm}^2$.

Magnetic permeability is defined in a similar manner to electrical permittivity and is a measure of the internal magnetic response of a material to an external magnetic field. Introduced by the wonderfully eccentric mathematical physicist Oliver Heaviside in 1885, the magnetic permeability links the external magnetic field $\mathbf{B}$ to the so-called auxiliary magnetic field $\mathbf{H}$ within the dielectric via the expression $\mathbf{H} = \mathbf{B}/\mu$. In SI units the value of the permeability of free-space is $\mu_0 = 1.2566370614 \dots \times 10^{-6} \text{ N A}^{-2}$.

The values of ϵ and μ will vary from one material to another and can also vary within a given material. According to the specific values these terms adopt, however, an impinging electromagnetic wave (to be described shortly) can be slowed down, sped up, and deflected. Within the metamaterials that enable invisibility cloaking devices the values of the permittivity and permeability are in fact negative (Fig. 4.10). Indeed, the art of constructing an invisibility cloak (as we shall see in Chap. 5) is all about engineering negative values of ϵ and μ for specific wavelengths of electromagnetic radiation – the very radiation, in fact, first described by a 33-year-old James Clerk Maxwell.

Making Waves

Maxwell's classic paper on the electromagnetic nature of light appeared in Volume 155 of the *Philosophical Transactions of the Royal Society* of London in 1865. The paper had been received on October 27, 1864, and read to the Fellows on December 8. The work is an intellectual *tour de force* with the somewhat unassuming title *A Dynamical*

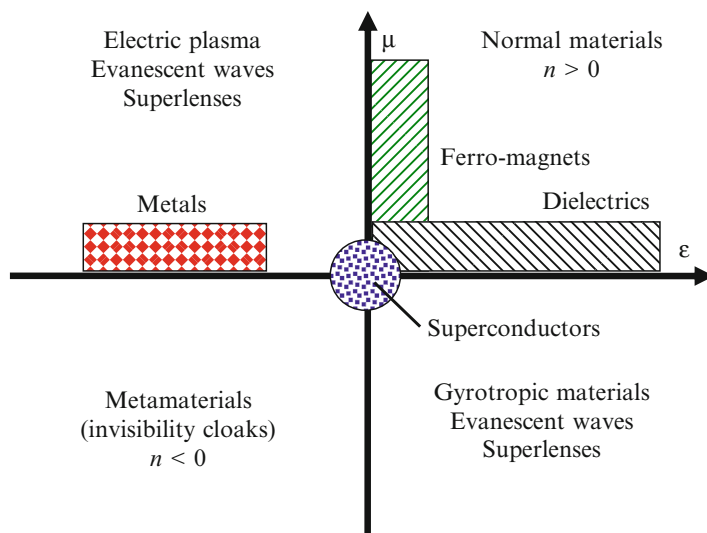


Fig. 4.10 Schematic diagram for the various positive and negative domains of electrical permittivity ϵ and magnetic permeability μ . Ordinary materials have positive values of both ϵ and μ ; metamaterials with a negative refractive index (the topic of Chap. 5) have negative values of both ϵ and μ . Superlenses (again described in Chap. 5) can be constructed from materials having a negative ϵ and positive μ . Materials with negative μ and positive ϵ are called gyrotropic and are responsible for such phenomena as the Faraday effect. Superconductors have very small values of ϵ and μ , and this enables electrical currents to flow with no energy loss to the carrying wire

Theory of Electromagnetic Fields. The paper was divided into seven sections, occupied some 53 journal pages, and is densely packed with mathematical equations.

After a short introductory preamble Maxwell explains the theory that he is about to propose can be “called a theory of the Electromagnetic Field, because it has to do with the space in the neighborhood of the electric or magnetic bodies, and it may be called a Dynamic Theory, because it assumes that in that space there is matter in motion, by which the observed electromagnetic phenomena are produced.” It is a field theory, therefore, and one that is imagined to be operating within an aetherial medium. Maxwell continues to explain in Sect. I of the paper that the aetherial medium must both pervade all space and be able to transmit at some finite speed electric and magnetic variations. Indeed, he developed a rather complex mechanical picture for the aether, requiring it to be composed of minute rotating cells (or vortices) separated by corridors of tiny spinning ‘idler spheres’ that carried the electric current (Fig. 4.11). The aether envisioned by Maxwell was capable of transmitting electromagnetic variations with no loss of energy and without becoming dispersed itself – a remarkable fluid, indeed.

Maxwell later explained that, “light is not a substance but a process going on in a substance.” We would now recast this in this form: “light is not a substance but a variation propagating through the local electric and magnetic field.” Maxwell’s suggestion for how the aether might transmit electric and magnetic field variations

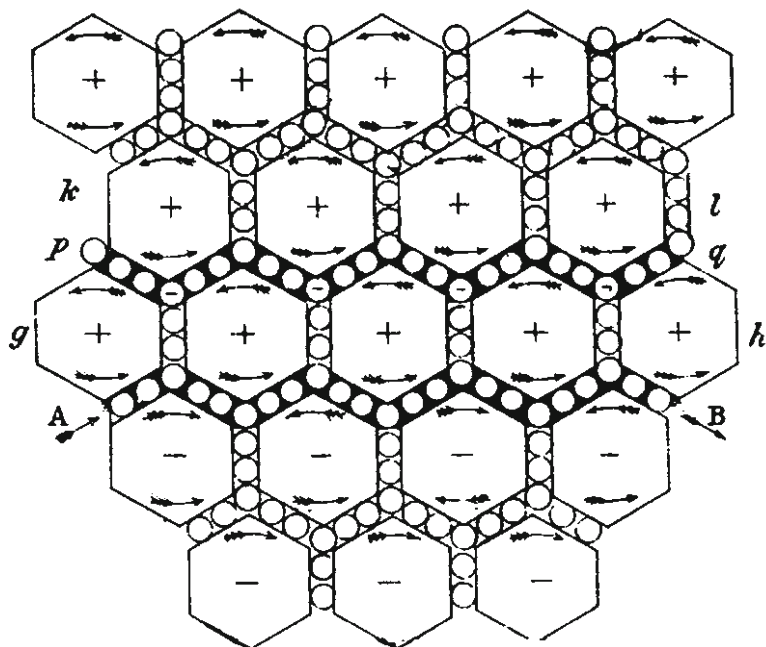


Fig. 4.11 Diagram from Maxwell's 1861 paper "On Physical Lines of Force." The rotating hexagonal cells are associated with magnetic field generation while the smaller circles rotate, conveying their rotary motion to the hexagonal cells, due to the passage of an electrical current. The picture is complex and actually contains a couple of graphical slips with respect to the arrow directions shown in the first and third cells in the last but one row

was a classic Victorian construct – it was after all the great era of industrialization and mechanics. The picture Maxwell presented was not necessarily intended to be taken literally – it was a model that could explain what might be going on. Luckily, perhaps, we need not think of space being pervaded by any strange brew fluid composed of invisibly small hexagonal cells and even smaller charge carrying spheres. The luminiferous aether, as envisioned by Maxwell, does not exist, but this is not to say that modern quantum physics hasn't produced an even more complex, even more bizarre picture of what the vacuum of space looks like. Even empty space is teeming with a veritable froth of virtual particles flashing briefly into and then out of existence.

Maxwell developed the electromagnetic properties of light in Part VI of his 1865 paper. He specifically considered the question that asks, if a wave is induced to move through a magnetic field, with what velocity will it propagate? Importantly, Maxwell finds that the speed of variation will be one divided by the square root of the product of the magnetic permeability (μ) and the electrical permittivity (ϵ) of the medium through which the wave is moving. The important point is that these constants are experimentally measurable and the quantity $1/\sqrt{\epsilon_0\mu_0}$ is, indeed, equal to the speed of light c in a vacuum. In a matter of fact, understated fashion,

Maxwell concludes that, “the velocity is so nearly that of light, that it seems we have strong reason to conclude that light itself is an electromagnetic disturbance.”

Not only is light an electromagnetic wave, it is now known to be a self-generating wave. This notion comes about since it is clear that light can propagate through the vacuum of space (indeed, starlight happily traverses trillions upon trillions of kilometers of it). The reason why light does not need a medium through which to propagate is embodied within the classical laws of Faraday and Ampere. In word form we can see the important symmetry between these two rules of nature:

Ampere's law: a time varying electric field produces an associated magnetic field
Faraday's law: a time varying magnetic field produces an associated electric field

If, therefore, some process initiates a variation in a magnetic field, then that variation will induce the formation of an electric field (Faraday's law), which will then induce the formation of a magnetic variation (Ampere's law), which in turn induces the formation of an electric variation and so on. In this self-creating manner, a sustained electromagnetic wave can be made to propagate through space without the need for any intervening medium. The other important and distinguishing characteristic of electromagnetic waves is that the electric and magnetic field variations are at right angles to each other (recall Fig. 2.12) and transverse to the direction of propagation.

In the spring of 1865 Maxwell resigned from his professorship at King's College London, retiring to the family estate in Glenair. There he quietly worked on developing his ideas relating to electromagnetism. In 1871, as a rather reluctant hero, Maxwell accepted the position of Cavendish Professor of Physics at Cambridge University. Once installed at Cambridge he helped establish what was to become one of the world's greatest laboratories for experimental physics. It was also while at Cambridge that Maxwell published in 1873 his now rightly famous work *Electricity and Magnetism*. With the publication of this final text, the end for Maxwell was rapidly approaching, and on November 5, 1879, Maxwell passed away at the relatively young age of 48 years, the doctor in attendance at this time commenting that, “no man ever met death more consciously or more calmly.”

Perhaps one of the worst criticisms that can be leveled against Maxwell upon developing his theory of electromagnetism is that he completely failed to see any practical use for it. This being said, no one else did either, until, that is, German experimental scientist Heinrich Hertz decided to perform a few simple experiments with a spark gap transmitter and a loop of wire in 1886 – and in the process discovered radio waves.

One of the vitally important elements of Maxwell's wave theory of light, as a form of electromagnetic radiation, is that the disturbance will always propagate at the same speed in a given medium. The theory, as such, makes no constraint upon the wavelength of the disturbance. In principle the waves can have any wavelength (recall Fig. 2.8). In the modern era the various forms of electromagnetic radiation that are known have wavelengths that vary from those of the ultra-short scale corresponding to gamma rays, which repeat over distances smaller than a trillionth of a meter, to radio waves, which vary on scales in excess of many meters. Gamma

rays, visible light and radio waves are all forms of electromagnetic radiation and all move through space at the speed of light. Maxwell, in fact, can be thought of as the great organizer; not only did he construct the first unified field theory (combining electric and magnetic fields), he also enabled the organization and unification of many diverse sets of experiments and observations.³

All Waves Together

The electromagnetic spectrum corresponds to the ordered run of electromagnetic waves according to increasing wavelength – or frequency. The shortest wavelength gamma rays have spatial variations smaller than 100 billionth of a meter, while the longest wavelength radio waves vary on scales in excess of many thousands of meters. Electromagnetic radiation with wavelengths between these two extremes are identified and named in Fig. 4.12. Visible light, that our eyes can see, falls in the wavelength region between 400 nm (blue light) and 700 nm (red light). Two scales are provided in Fig. 4.12: frequency, f , and wavelength, λ , and for all electromagnetic waves the product of these two quantities is equal to the speed of light.

The electromagnetic spectrum covers a tremendous range of wavelengths, and this is the ultimate technical challenge to those engineers who might attempt to make a universal invisibility cloaking device. Indeed, no one specific metamaterial will correctly channel all wavelengths of electromagnetic radiation. While cloaking over all wavelengths is a Type II impossibility it is worth noting that there are currently no specific reasons for wanting a universal invisibility cloak anyway. From a military perspective, for example, cloaking at visible, infrared and radar wavelengths are the main domains of interest.

Lights Out for the Aether

While Maxwell had successfully described the characteristics of electromagnetic radiation, there was still the mystery of the aether to sort out. Fortunately, it was a mystery that ultimately vanished – literally. The Michelson-Morley experiment was all about precision, and if the speed of light varied by just the smallest of fractions with respect to the direction of travel through the aether, then this experiment would reveal it. The idea is that as Earth orbits the Sun over the course of 1 year it will be alternately moving with, and then against, 6 months later, the direction of the aether's

³ Paul Davies (Arizona State University) in his recent book on the search for extraterrestrial intelligence, *The Eerie Silence* (Mariner Books, 2010), has suggested that if we want to signal our intelligence to ET then we should broadcast Maxwell's equations. This, he argues, would be a clear indication of our advanced scientific development and prowess.

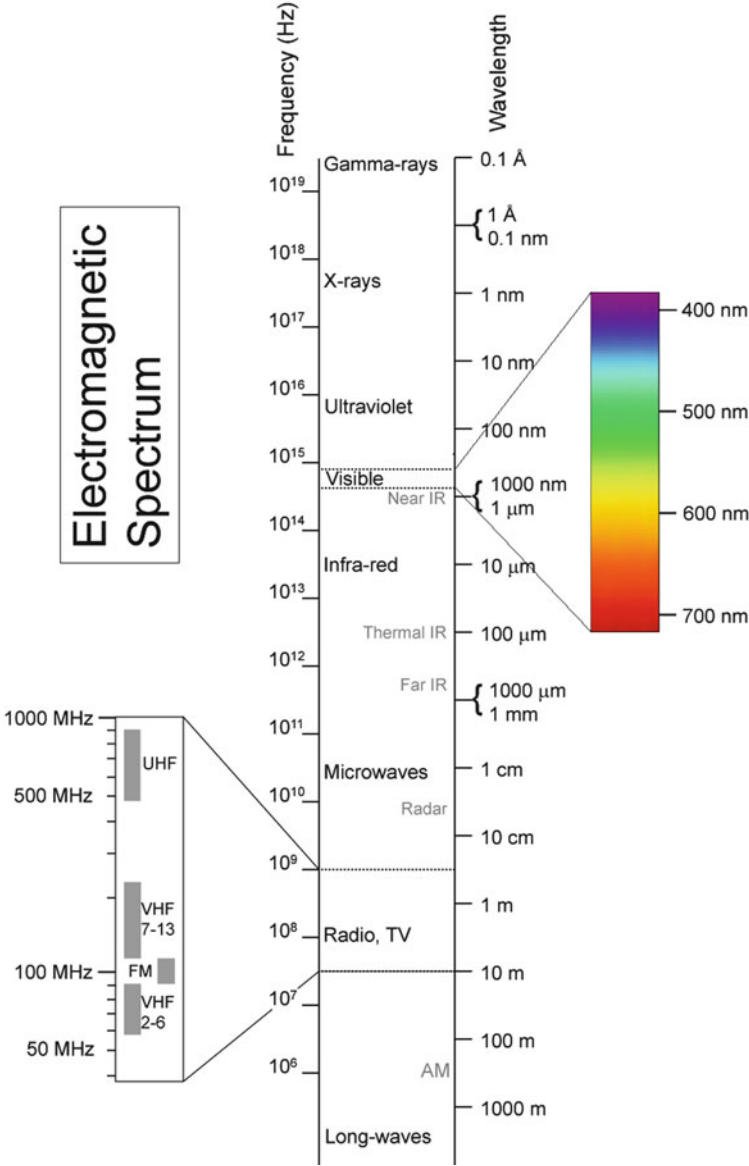


Fig. 4.12 The arrangement of the electromagnetic spectrum. Visual light occupies just a small segment of the wavelength region between γ -rays and radio waves (Image courtesy of Wikipedia commons.)

orientation (Fig. 4.13a). The difference in these respective motions, everyone (including Michelson and Morley) believed, should be measurable (Fig. 4.13b).

Michelson and Morley performed their now classic experiment in 1887, but to their surprise found no noticeable variations. This, however, was a glorious and

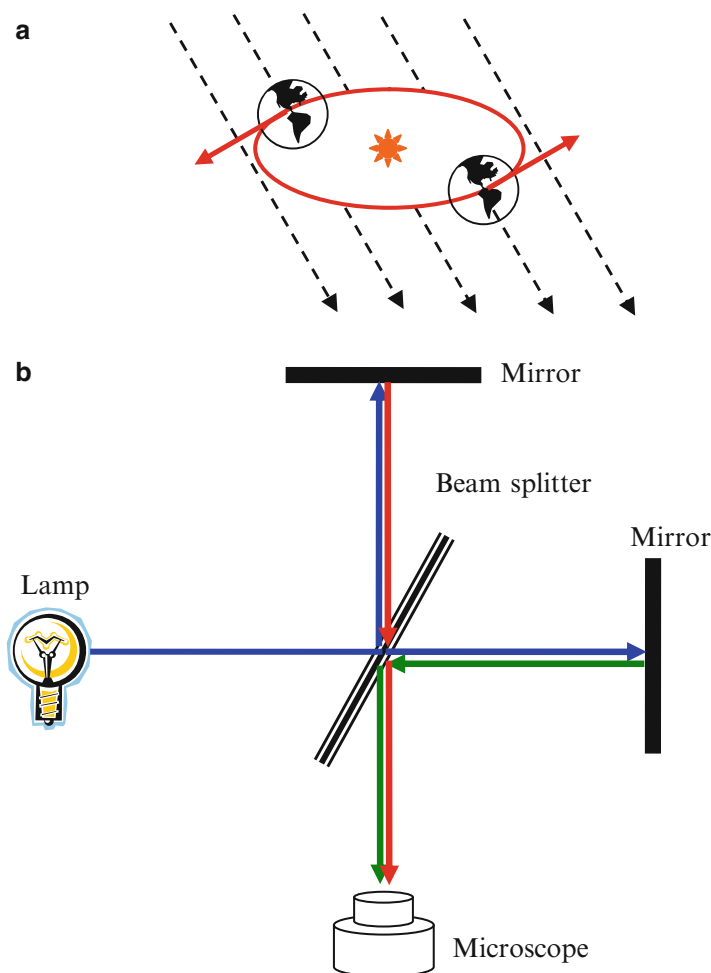


Fig. 4.13 (a) Earth's motion through the aether should produce a cyclical 6 monthly variation in the measured speed of light. The dashed lines indicate the direction of the Solar System's motion through the aether, while the solid lines indicate the direction of Earth's motion in its orbit about the Sun. (b) Schematic of the Michelson–Morley interferometer. In this experiment a single light beam is sent along two paths at right angles to each other via a beam splitter. If the speed along one path was different from the speed along the other then an interference pattern (recall Fig. 2.9) would be observed through the microscope eyepiece

revolutionary null result – perhaps the most famous non-detection in the history of science? The measured speed of light showed no variation with respect to the time of year (that is, with Earth's location in its orbit) nor upon the orientation of the experiment in the laboratory. The conclusion was straightforward (well, to some), and it was clear that there was no such entity as the aether pervading space.

The physics community of the day was not overly perturbed by the outcome of Michelson-Morley's experiment. Certainly, the result cast the idea of the aether into clear doubt, but in the pragmatic fashion of science it was typically assumed that further searching and experimentation would ultimately explain what was going on. What was going on, however, would only become clear in the new twentieth century, in 1905 in fact, when a young Albert Einstein, an instant superstar, burst onto the scientific scene.

Putting the Particles Back

Question: When is light not a wave?

Answer: When it is a photon.

The double slits experiment by Thomas Young (Fig. 2.10) had supposedly made the particle theory of light redundant, but rumors of its demise turned out to be somewhat premature. Indeed, the particle theory of light was to come back with a vengeance, ushered in at the beginning of the twentieth century, riding high on the wave that was the new physics of quantum mechanics.

Albert Einstein became less and less comfortable with the developments in quantum theory as time went by, but ironically he was responsible for establishing it as a serious field of study. While German physicist Max Planck had introduced the idea of quantized energy in 1900, he did so as a mathematical maneuver rather than one relating to physical understanding. It was Einstein in 1905 that first took Planck's ideas seriously and used them as a means of explaining a strange phenomenon that had, in fact, been described by Heinrich Hertz in 1887. Hertz noted that when certain types of metals were exposed to certain frequencies of light they begin to emit electrons. Specifically, and importantly, there appeared to be a threshold frequency below which there was no effect, but above which the electrons began to be emitted, with the energy of the ejected electrons increasing with the increasing frequency of incident light. What was strange about this phenomenon was that according to the wave theory the light, the energy of the ejected electrons, should increase with the wave intensity, not the frequency. Einstein solved this problem in terms of what he called light quanta – now known as photons.⁴

Einstein's 1905 result was a watershed moment in physics, and resulted in the then new concept of the dualistic nature of light – more often now referred to as the wave-particle duality of light. What this concept means is that light has both wave-like and particle-like characteristics – it is a hybrid entity, being neither a true wave nor a true particle, but something that is simultaneously both of these phenomena and more. Indeed, according to modern concepts we shouldn't actually think of photons as having either size (see, however, Appendix 3) or mass. Light, in some very strange sense and in spite of being everywhere, has no corporeal presence.

⁴ The term photon was actually coined by American chemist Gilbert Newton Lewis in 1926.

Absolute Light Speed

The year 1905 was a pretty good one for Einstein. Not only did he produce his Nobel Prize-winning work on the photoelectric effect, he also published his first paper on the special theory of relativity. This revolutionary work and re-interpretation of motion at very high speeds was a direct result of the 1887 Michelson and Morley experiment. The experiment, you might recall, showed that there is no aether and no preferred frame of reference with respect to light propagation. Einstein realized that this observation implies that all observers, no matter where they are in the universe and no matter how they are moving will always find that light travels at the same speed. Indeed, Einstein argued that the speed of light is an absolute limit, and this, in fact, is the underlying principle of his theory of special relativity.

In the everyday world we are used to the idea that space and time are absolutes. Time ticks away at a constant rate, and the distance between any two fixed objects is always the same. Likewise, when something is moving in our everyday world, its motion is always measured relative to how fast we ourselves are moving. At extremely high speeds, however, it is speed that becomes absolute – i.e., that of the speed of light – and accordingly space and time become relative. To accommodate this limiting speed, Einstein showed that clocks will tick at different rates and sizes will vary according to different observers when the speeds concerned are relativistic. To our everyday take on the world these concepts seem ludicrous, but they have long been shown to be real and observable effects in countless numbers of experiments.

Warp Speed and Absolute Invisibility Cloaking

Einstein enters our discussion one more time. Moving beyond his pioneering studies of the photoelectric effect and special relativity, Einstein introduced to the world his theory of general relativity in 1915. This more expansive theory offered a way of accounting for gravity in terms of the curvature of four-dimensional spacetime – as we saw in Chap. 3.

The world of general relativity is a mathematically complex one, but for our purposes all we need to know is that light travels along the shortest possible path within spacetime, and space curves according to the concentration of mass at a specific location. Since all forms of electromagnetic radiation, from gamma rays to radio waves, must move within space, so the ultimate invisibility device would be one in which the spacetime around an object (say a spacecraft) is manipulated in such a way that electromagnetic radiation simply flows around it. This all sounds too good to be true – and, of course, it is. The process of turning the curvature of spacetime on and off at command, while possible in principle (that is, it is possible according to a series of mathematical equations) it is most likely impossible in practice. As to whether we are dealing with a Type II or Type III impossibility is presently unclear and partly depends upon the existence, or not, of variously proposed exotic

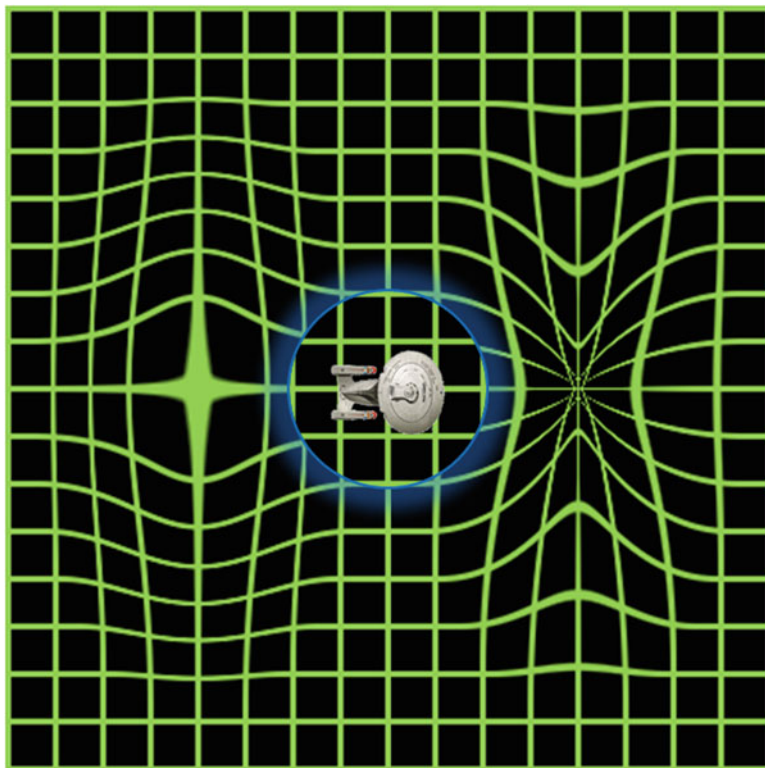


Fig. 4.14 Visualization of a warp field drive, with a spaceship surrounded by a bubble of normal spacetime, being carried along between two regions of contracted (ahead of the ship) and expanded (behind the ship) spacetime (Image courtesy of *Wikipedia* commons)

forms of matter. If nothing else, the idea of an absolute invisibility cloak sets a boundary upon what might be imagined and eventually achieved.

In spite of the uncertain, if not insurmountable technical difficulties, physicists Miguel Alcubierre (National Autonomous University of Mexico) has studied a similar problem to the warped space, all wavelengths, invisibility cloak. Specifically, Alcubierre asked if it might not be possible to warp the space around a spacecraft as a means of providing propulsion. The key step, he argues, is to find a way of expanding space behind the spaceship and contracting space in front of it (Fig. 4.14). This might seem like a mad idea, but that is no reason to reject it out of hand. Indeed, as we shall see in the next chapter the newly emerging technology of metamaterials might even enable laboratory-based testing of Alcubierre's warp-drive concept.

While the realization of practical warp-fields lies (infinitely?) far off in the myopic future, we do have at this point a better, albeit somewhat complicated, picture of what light actually is – or, perhaps more to the point, what it isn't. Indeed, in this chapter light has been revealed as something more than a wave and more than a

particle. Moreover, it is a phenomenon that requires no medium through which to propagate, and it is composed of time-varying magnetic and electric fields. Taking all these ideas on-board, in the next chapter we shall see how the propagation of light can be controlled by manipulating the magnetic and electrical properties of the medium through which it is moving – indeed, we shall see how light can be made to travel in ways that nature, that great experimenter, has never physically realized.

Chapter 5

The Metamaterial Revolution

“Everything not forbidden is compulsory.” So writes T. H. White in his book *The Once and Future King* – and, if nature ever had need of a motto then surely it would be these words. For indeed, nature is the ultimate utilitarian, ever building, ever moving, never stationary, searching in all possible ways, step by small step, for a better adapted, more efficient, more successful survivor. If a process is not physically impossible, that is, against the laws of physics, nature will usually find a use for it.

Metamaterials, however, appear to be one trick that nature has missed – so far at least. More than just combining new chemical elements or various metals to make new molecules and alloys, metamaterials are custom substrates built from the bottom up, with very specific layering and shaped inclusions. Indeed, the word *metamaterial* is the concatenation of *meta*, meaning “beyond,” and *material*, meaning, “natural,” or in this case, “manmade substructure.”

With respect to invisibility cloaking, the key requirement of the metamaterial substrate is to interact with and carefully control the passage of electromagnetic radiation through it – the light needs to be seamlessly channeled around a cavity in which a solid object can be hidden from external view. Crack the manufacturing technology for metamaterials and you will soon have a fully functional invisibility cloak.

The Black Arts

In an ideal world it might be possible to make an aircraft, missile system or building fully invisible, at all wavelengths of light, but the question arises as to the practicality of such an approach to cloaking – the point being that a military aircraft, for example, need only be invisible, or at least difficult to detect, at those wavelengths in which the enemy is monitoring its airspace. In practical terms this means that an aircraft need only be invisible at radar wavelengths. A building or ground-based missile system, in contrast, might be best cloaked at optical wavelengths so that it is difficult to photograph and/or identify correctly – no one will look for a building using X-rays, so it is pointless to cloak it at such wavelengths.



Fig. 5.1 The Lockheed F-117 Night Hawk stealth bomber. Used in active service by the USAF from 1981 to 2008

Detection and recognition of military targets typically proceeds by both active and passive means (recall Chap. 2). Radar detection, for example, is an active detection process in the sense that a radar signal is sent out by a transmitting antenna, and the object is detected by the fact that it reflects some of the radar signal back to a receiving antenna (recall Fig. 2.1). In contrast, looking for an aircraft or missile in flight in the infrared part of the spectrum is a passive process. In this latter case one is monitoring the actual heat emission from an object rather than looking for a returned, that is, reflected signal. Modern military aircraft and missiles are accordingly designed to minimize both their heat signature and their so-called radar cross-section. It is by considering these design issues that the seemingly odd shape of the stealth bomber (Fig. 5.1) came about – indeed, these aircraft are designed more for minimizing their detection than their flying characteristics.

The designers of the stealth bomber have utilized many different approaches to make it less readily detectable. Firstly, the odd-looking, angular flat panels on its front surface are used to reflect radar signals in a direction away from that of the incident beam. In other words the radar signal is not reflected directly back towards the receiving equipment but straight upwards into the atmosphere; the basic design mantra being, “No returned signal means no aircraft detection.”

Second, the engines of the stealth bomber are set back deep within the aircraft’s mono-wing. This is done to minimize the heat signature of the engines, and it also reduces the turbulent wake trailing behind the aircraft – a feature that can be detected and thereby used to betray the existence of the aircraft to enemy reconnaissance.

Thirdly, the surface of the aircraft is coated with specially formulated radar absorbing materials and paint. Again, the philosophy here is that if the incident radar signal is absorbed at the aircraft surface then it isn't being reflected back to the monitoring receiver, and the aircraft goes undetected.

The specialized surface paint, tapes and reflectors, however, have turned out to be the Achilles heel of the stealth bomber. The aircraft can only fly under certain weather conditions, and the surfaces must be maintained in pristine condition at all times.

It has been a classical give and take, evolutionary kind of battle, and ever since active radar detection was first introduced during World War II engineers have struggled to find new and ever more inventive ways of making aircraft more difficult to find and target. To this end numerous radar absorbent paints and special airframe structures have been developed. Radar absorbent, also called iron ball, paint, for example, contains small spheres coated with carbonyl iron or ferrite (a very pure grade of iron invented by Germany chemical company BSAF in 1925). The idea here is that rather than being reflected, the energy of any incident radar beam is absorbed by the molecules in the coating applied to the spheres, causing them to heat up. The heat is then conducted into the interior of the aircraft and harmlessly dissipated. In some sense radar absorbing paint is the antithesis of retroreflecting paint (recall Chap. 3), fully absorbing the incident radiation rather than fully (or nearly so) reflecting it.

Another way in which an aircraft's radar signature can be reduced is to force the radar wave to destroy itself, a process that can be achieved by destructive interference (recall Fig. 2.9). In this case the idea is to have a multilayered outer skin with the components separated by one-quarter wavelength of the radar signal. The wave reflected from the first surface will then destructively interfere with the wave reflected from the second, back surface, since it will be one-half of a wavelength out of phase. This is all well and good, of course, if you know the wavelength of the radar that the enemy might be using. Ideally what is needed is a broad wavelength range of visible, infrared and radar absorbing surfaces and paints, and this is where the new metamaterials come into their own – they are the designer, mix-and-match, pick-your-own version of camouflage invisibility cloaks.

But how did all this begin? The story in fact starts with an abstract idea, which then lay totally ignored for nearly three decades.

The Strange Paper by Dr. Veselago

Sometimes, rarely but sometimes, nature can be bested by human ingenuity. One such example of this occurred in a research paper published¹ in January of 1968 by Victor Veselago of the Moscow Institute of Physics and Technology, Russia (Fig. 5.2).

¹ The English translation of Veselago's paper incorrectly states that the original Russian language article was published in 1964 – it was actually published in 1967. This being said, some of the effects relating to negative refractive index can be traced back to Russian papers published in the mid-1940s and 1950s.

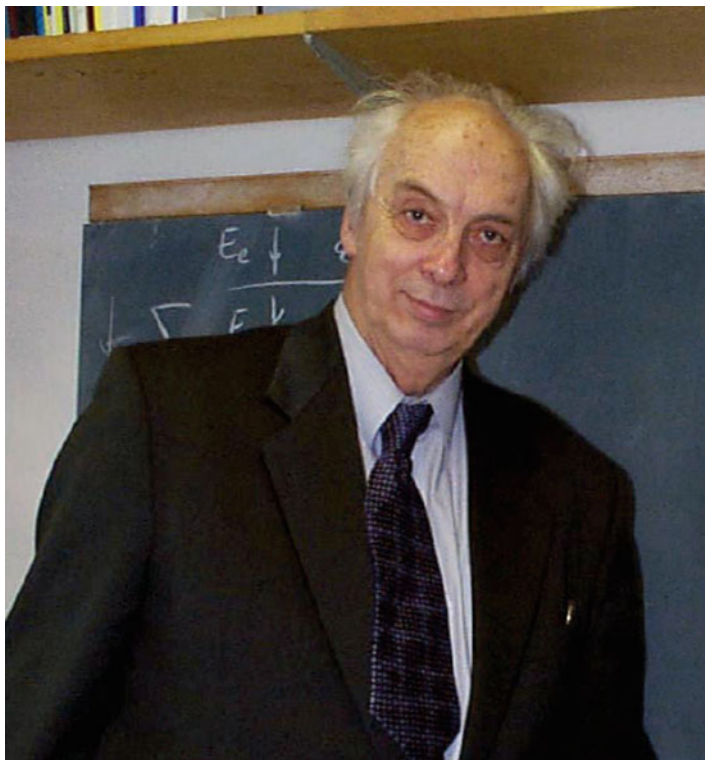


Fig. 5.2 Victor Veselago (b. 1929) – early pioneer of negative refractive index research (Image courtesy of Wikipedia commons)

This paper, under the all-revealing title of *The Electrodynamics of Substances with Simultaneously Negative Values of ϵ and μ* , was published in the journal *Soviet Physics Uspekhi*. The paper is just six pages long and considers what, at first glance, is a seemingly mad idea, but an idea, nonetheless, that is mathematically allowable. Indeed, it is a mathematical trick that nature, as far as is known, has managed to miss. While the initial response to the paper was far from enthusiastic, and it was essentially ignored as an interesting but quirky idea for 30 years, it eventually resulted in the investigation and construction of metamaterials – materials that have no counterpart in the natural world. Indeed, they are the ultimate in manmade constructs. At the heart of metamaterial design, however, is the manipulation of the magnetic permeability and electrical permittivity.

In the last chapter we saw that one of Maxwell's most important clues as to the nature of light was that electromagnetic waves traveled at the speed of light. This was given explicitly in the sense that the wave velocity is $V = 1 / \sqrt{\epsilon\mu}$ where ϵ is the electrical permittivity of the medium and μ is the magnetic permeability. The next step, the daring one taken by Veselago, is actually a simple mathematical one.

The genius of Veselago, however, was to see that the mathematical trick could make physical sense. So, what is this trick? It is literally as simple as $1 \times 1 = 1$, or more to the point it is equivalent to $-1 \times -1 = 1$ – this is just the mathematical rule that dictates the product of two negative numbers is always a positive number. Now, of course, Veselago wasn't just playing with the multiplication of 1 or -1 , which is neither mathematically nor physically very interesting, but what he was looking at is the square root term in the velocity equation for an electromagnetic wave: $\sqrt{\epsilon\mu}$. In real number² mathematics only the square roots of positive numbers are allowed, and accordingly the product $\epsilon\mu$ must be positive. Now, we saw earlier, in Fig. 4.10, that both the electrical permittivity and the magnetic permeability are generally positive terms and accordingly, their product must also be positive. OK, all well and good. But now, with our mathematical trick in mind, we also know that the product $-\epsilon \times -\mu$ must also be positive, but this directly leads to the question – the one asked by Veselago – what physical sense can be attributed to the negative values of ϵ and μ ? Indeed, it was by considering this question that the idea for metamaterials and the notion of negative refractive index first came about.

Veselago's article is a classical physics paper in the sense that the topic was being considered for purely intellectual reasons. In the opening paragraphs Veselago explains that while no substances in nature are known to have simultaneously negative values of permittivity and permeability, one can at least work out, "where and how one can look for such substances. Since in our opinion the electrodynamics of substances with $\epsilon < 0$ and $\mu < 0$ is undoubtedly of interest, independently of our no[t] having such substances available, we shall at first consider the matter purely formally." Veselago's analysis began, as indeed it must, by assembling together Maxwell's equations. Now, however, the negative values of ϵ and μ are inserted into the mix, and the consequences for the wave equation are analyzed. Keeping careful account of all the negative and positive terms, Veselago found that while the velocity of the wave is still related to the square root of the product $\epsilon\mu$, the solution must now have a negative sign in front of the square root symbol. Indeed, the refractive index is negative: $n = -\sqrt{\epsilon\mu}$. To gain some idea of how important Veselago's paper has become we can look at the number of citations it has received in published research papers since it first appeared in 1968. Using the NASA-administered Abstract Data Service (ADS) to determine the yearly counts Fig. 5.3 shows the cumulative sum of references in the time interval from 1968 to 2010. In total there are 1,267 references to Veselago's paper, and the rapid rise in interest in Veselago's work, following the publication of John Pendry's article in *Physics Review Letters* in 2000 (to be described shortly) is clearly evident.

²Mathematicians divide number domains into the integers, the reals and the complex. The integers are all the positive and negative whole numbers (including zero), whereas the real numbers include all the fractional and integer numbers between minus and plus infinity. The complex numbers are composites of a real number and an imaginary term that is multiplied by i , defined as being the square root of -1 (hence imaginary number). Remarkably, many physics and engineering calculations are cast in terms of complex numbers. This latter result, as physicist Roger Penrose has succinctly put it, "is a cause for wonder."

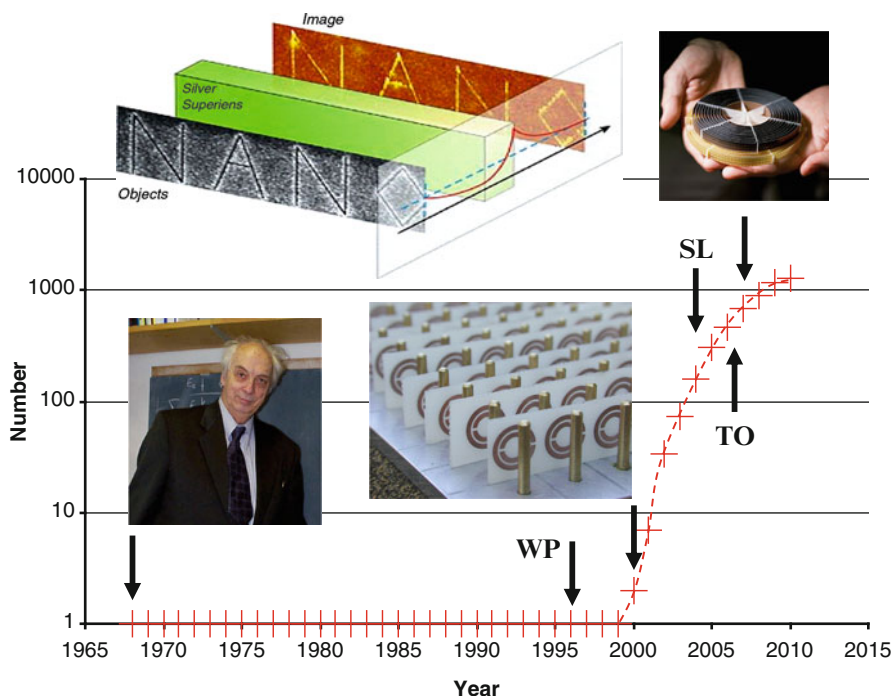


Fig. 5.3 The cumulative total of papers per year (crosses) making reference to Victor Veselago's 1968 paper on negative refractive index. The times at which the first negative refractive index metamaterial superlens (SL) and microwave cloaking devices were announced are indicated with arrows (all these topics will be discussed later). WP corresponds to the year in which the paper responsible for stimulating much of the present-day interest in metamaterials by Andrew Ward and John Pendry was published. TO refers to the publication date of the now classic papers on transformation objects and invisibility devices by Ulf Leonhardt and John Pendry and colleagues at Imperial College, London. The image in the upper left-hand corner is the schematic for a super-resolution, metamaterial lens – the first such devices being developed during 2004 and 2005

Once interest took off, the number of references to Veselago's paper was initially doubling every 6 months, but in recent years this has leveled off to a doubling rate of every 2 years. This latter slowing down probably has more to do with the specialization of the research field, rather than any lack of appreciation for Veselago's contributions. The most recent research publications are essentially building upon results obtained within just the last few years. Nonetheless, after 32 years of being ignored, the rise in interest, once it did begin, in Veselago's paper was remarkable and meteoric.

Negative refractive index does not change Snell's law, but it does change the refraction rule (as described in Chap. 2). When a light ray encounters a negative refractive index surface, rather than being refracted at an angle to the right of the surface normal, it is refracted to the left (Figs. 5.4 and 5.5). It is for this reason that negative refractive index metamaterials are sometimes referred to as left-handed

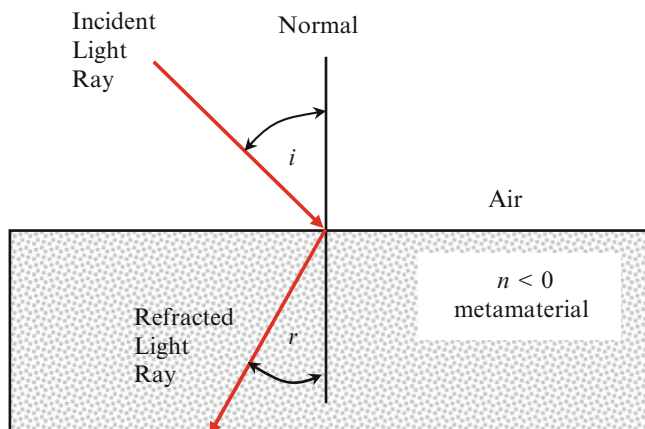


Fig. 5.4 Refraction of a light ray upon encountering a negative refractive index metamaterial. This diagram is in contrast to refraction in normal matter as indicated in Fig. 2.3b



Fig. 5.5 Computer generated images showing the refraction expected for a glass rod placed in (left) a glass containing a negative refractive index fluid ($n = -1.33$) and (right) a glass containing water ($n = 1.33$) (Image courtesy of The Dionne Group, Stanford University)

materials. In essence, the light ray is refracted in the opposite way to (normal) positive refractive index materials (which are accordingly called right-handed materials). Indeed, one might say that negative refractive index materials open up a mirror-world of refraction optics.

Having explained the consequences of light encountering a negative refractive index medium, Veselago looked even deeper into the equations and found that

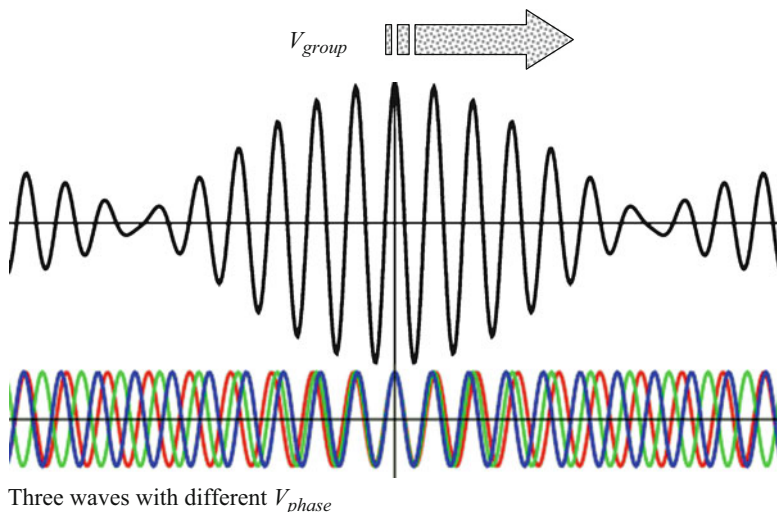


Fig. 5.6 The formation of a wave packet in a dispersive medium. The group velocity corresponds to the speed with which the wave packet itself moves through the medium. The wave packet shown here (*top*) is the sum of the lower three phase velocity curves (Image adapted from E. L. Wright, UCLA)

additional odd features should appear. Perhaps the oddest result that he found relates to the way in which an electromagnetic wave, and its associated energy, will propagate through a left-handed substrate. With respect to the propagation of an electromagnetic wave through a non-conducting medium there are two properties that need to be considered, the so-called phase velocity and the group velocity. The phase velocity is related to the actual wavelength and frequency of the incident light, with $V = \lambda f$ (recall the discussion from Chap. 2). In a vacuum the phase velocity is the same for all wavelengths, but in a so-called dispersive material (such as a metamaterial) the speed will vary slightly according to wavelength. In this way an interference pattern (recall Fig. 2.9) will be set up and a modulating wave profile, or wave packet, will form. The so-called, group velocity gives the speed with which the profile of the wave packet moves through the medium (Fig. 5.6).

In an ordinary, right-handed, material the individual waves, the wave profile and the energy carried by the wave all propagate in the same direction—although the phase and group velocities need not have the same numerical value. In a negative refractive index, left-handed, medium, however, it turns out that the individual waves move in the opposite direction to the wave packet profile and the energy. The phase and the group velocity have different signs and move in opposite directions from each other. Remarkably, we have the situation that the individual waves (that interfere to make up the wave packet) travel backwards, while the wave packet itself and the energy associated with it move forward into the medium. Indeed, it is exactly for this reason that the refractive index is negative. John Pendry has described the strangeness of this situation by noting that if a flashlight emitting a continuous beam

of light was embedded in an optical metamaterial with a negative refractive index then if one could follow the individual waves they would appear to be moving backward from the metamaterial into the flashlight, as if one were seeing a movie running backwards. The energy and wave packet, however, travel away from the flashlight into the metamaterial. Left-handed metamaterials essentially force us to re-write the book of optics and electromagnetic radiation propagation.

An additional oddity associated with waves moving within left-handed metamaterials is that the Doppler effect will behave in the opposite manner to that normally experienced. This effect is most commonly encountered, in every day life, through sound waves, with an increase followed by a decrease in the pitch of a siren on an emergency response vehicle, for example, as it moves towards and then away from you. In a left-handed metamaterial substrate this phenomenon is reversed, with a decreasing pitch as the sound source moves towards you and an increase in pitch as it moves away. The inverse Doppler effect, at radio wavelengths, was first demonstrated in a negative refractive medium in 2003 by Nigel Sneddon and Trevor Bearpark at the BAE Advanced Technology Center in Bristol, England. The inverse Doppler effect for light was first demonstrated in a metamaterial by Jiabi Chen (University of Shanghai for Science and Technology) and co-workers in 2010. Left-handed metamaterials really do behave in the strange manner that Victor Veselago said they would.

Marconi All at Sea

If Victor Veselago can be credited with the first detailed investigation of materials having a negative refractive index, so it was Sir John Pendry, professor of theoretical physics at Imperial College, London, who brought the work into the main stream of science (Fig. 5.7). The story of this journey, however, is a long and convoluted one, containing many twists and turns.

It all began with the British Navy – well, more to the point with one of its key technical service contractors, Marconi Materials Technology. As part of their research into minimizing the radar signature of naval warships the Marconi engineers had developed a new carbon material with all the required characteristics. The problem was, however, they had no clear idea as to why it worked. It was at this point, then in the mid-1990s, that Pendry, as a renowned researcher in solid state physics and electromagnetic theory, was called in. Could he find out how the material functioned, they asked?

The answer, Pendry discovered, lay within the fuzzy fiber strands out of which the radar absorbing carbon material was made. These strands, it transpired, were just the right length and shape to interact with the electrical component of a wave with a wavelength in the radar region of the electromagnetic spectrum. The various strands in the fiber essentially acted as small receiving antennas, and when the radar beam fell upon them they were able to absorb part of its energy, thereby reducing the strength of any reflected wave. The reduced strength of any returned radar signal, of course, was exactly what the Navy was interested in.



Fig. 5.7 Sir John Pendry FRS (b. 1943), whose groundbreaking theoretical work inspired the construction of the first invisibility cloaks and superlenses (Image courtesy of Imperial College)

Having sorted out the theoretical underpinning of the radar absorbing material, Marconi MT suggested that Pendry might try to do some further investigations. One topic that surfaced and seemed particularly interesting was that relating to material magnetism. Might it be possible, the reasoning went, to develop a material that could have its magnetic characteristics turned on and off according to the electromagnetic radiation incident upon its surface? This, Pendry realized, was a whole new area of material science research—and a highly interesting and potentially

profitable one, too. A plan of research eventually emerged. The aim was to see if, in essence, an artificial magnet could be made out of material that was not of itself magnetic.

As we saw in Chap. 4 magnetism is all about charges moving around in loops, and so Pendry and co-workers envisioned making a material that had as part of its internal structure numerous small copper wire loops. Although copper is a conducting metal it is not a naturally magnetic one, but, the key point is, if one can induce electrons to circle around a loop made of copper wire then according to Ampere's law a magnetic field will be produced. The secret to making such a nonmagnetic material come to life, therefore, lay in the activation of its internal copper loops. Further, by changing the properties of the copper loops (technically known as split ring resonators—to be discussed below) then the magnetic properties of the metamaterial could also be changed. The Imperial College research group published their results in 1999, and a veritable firestorm of activity erupted thereafter (recall Fig. 5.3).

Negative Refractive Index Engineered

Experimental physicist David Smith at the University of California in San Diego was instantly hooked on the topic of metamaterials as soon as he heard about Pendry's ideas, and he soon realized that experimental verification was possible. He mulled over ideas with colleague Willie Padilla, and a working negative refractive index metamaterial model for microwaves began to be pieced together in their laboratory. The technical write-up of the group's research was published in the May 1, 2000, issue of the journal *Physical Review Letters*. The title of the paper was short and to the point: *Composite Medium with Simultaneous Negative Permeability and Permittivity*. Within its brief four pages of text a theoretical model for the construction of a negative refractive index metamaterial (Fig. 5.8) is described along with the details of how its properties can be experimentally verified. The crucial experiment related to how microwaves with slightly different wavelengths would propagate through the medium – it was a transmission mode variation that they were looking for. Importantly, the way in which the waves moved through the metamaterial were exactly as all the computer modeling predicted. Light, at microwave wavelengths, in the first-ever purpose built negative refractive index metamaterial was behaving as it was supposed to. The paper concluded triumphantly, “We are now in a position to further investigate the fascinating electrodynamic effects anticipated for...composite metamaterials.”

The next experimental step was to verify the negative refraction properties of their metamaterial. This Smith and co-workers soon achieved, publishing their results in the journal *Science* on April 6, 2001. A new more compact metamaterial substrate was fabricated for these follow-on experiments, and now the plan was to measure the angle through which the incoming microwave was refracted after encountering the metamaterial lattice.



Fig. 5.8 The first purpose-built metamaterial to test the ideas of negative refractive index media as theoretically described by Pendry et al., and Vesalago. Copper loops (split ring resonators) have been etched on to fiberglass circuit board substrates, and the wire and resonator substrates have been arranged into a lattice structure (Image courtesy of University of California at San Diego)

By setting their metamaterial substrate at an angle to the incident microwave beam (Fig. 5.9), Smith and co-workers could measure the corresponding angle of refraction, and accordingly they found $n = -2.7$; the refractive index was indeed negative. With the publication of the Smith et al., paper in 2001 the verdict was in; negative refractive index metamaterials could be constructed, and they behaved in exactly the way that Victor Vesalago and John Pendry and the Imperial College research group said they should.

Split Ring Resonators and Metamaterials

At the heart of the first fabricated negative refractive index metamaterials is the split ring resonator (SRR). Indeed, John Pendry has described SRRs as, “the metamaterial equivalent of a magnetic atom.” It is this structure that interacts with, responds to and then back-reacts upon the impinging electromagnetic wave. An SRR is composed of two nested copper C-shaped rings. When an electromagnetic wave encounters an SRR its oscillating magnetic field causes electrons in the SRR to move in sympathy with it, thus enabling an electric current to flow. The gap in the ring component acts like a capacitor, briefly storing a charge before letting it proceed, and this establishes

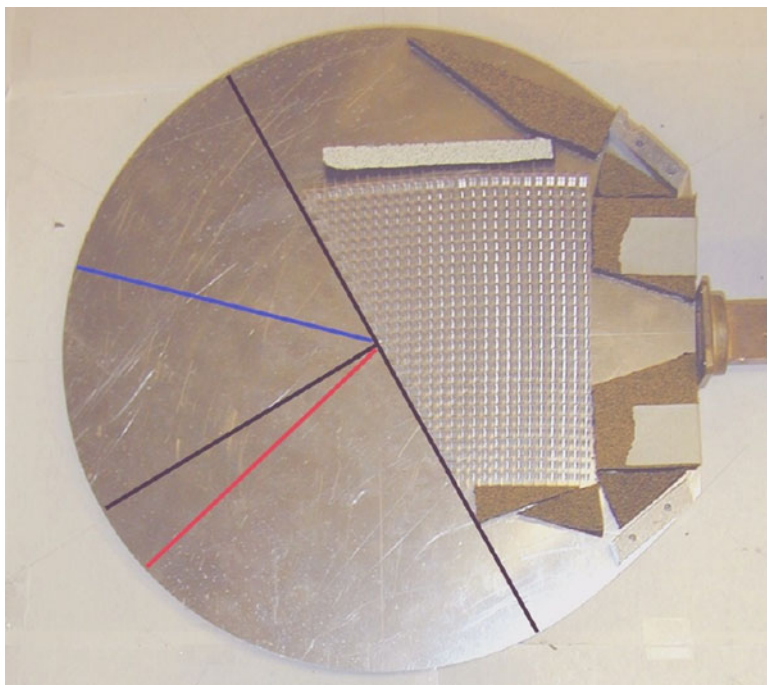


Fig. 5.9 The experimental arrangement for determining the refractive index of the metamaterial substrate shown in Fig. 5.8. The surface plane and normal are shown by the *black* lines. Microwaves enter from the right and are refracted in the direction of the *red* line, indicating the negative refractive index of the metamaterial. The *blue* line indicates the direction the ray would take if the refractive index were positive (Image courtesy of University of California, San Diego)

an oscillating electric field that back-reacts upon the impinging electromagnetic wave. The wavelength range over which the magnetic permeability μ is negative is controlled by the radius of the inner C-ring as well as the separation between the inner and outer C-rings. The electrical permittivity ϵ is further determined by an array of straight conducting wires. In this case, the electrical field of the incident wave sets up an oscillating dipole magnet in each wire, and the electrical permittivity is controlled by adjusting the wire spacing and length. It is the combination of the linear wire and SRR arrays (Fig. 5.10), along with their relative size and spacing that determines the wavelength region over which both ϵ and μ will be negative.

A metamaterial is fabricated by engineering an arrangement of repeated SRRs and vertical wire cells. Provided each cell and the spacing between adjacent cells is much smaller than the wavelength of the incident electromagnetic radiation, then the metamaterial will act like a homogeneous material – in this sense, the individual cells within the metamaterial combine to work like a kind of crystal made of artificial atoms (the SRRs). The crystal-like analogy is nicely illustrated by the so-called Boeing Cube (Fig. 5.11), which was designed by Minas Tanielian and co-workers in 2004 at Boeing Phantom Works in Seattle. Indeed, the Boeing team was working

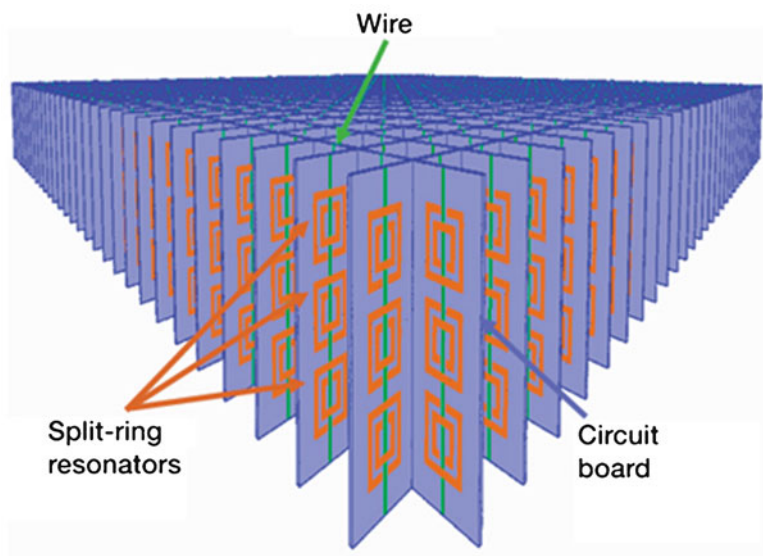


Fig. 5.10 Schematic of a metamaterial composed of SRRs and a parallel wire grid array (Image courtesy of Berkeley Labs)

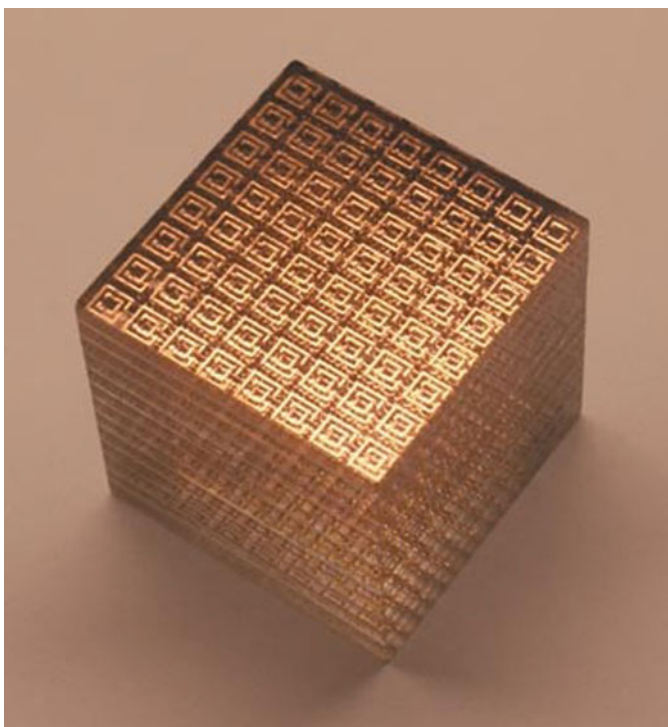


Fig. 5.11 The Boeing Cube. The component copper SRRs cells and wire array are visible in this metamaterial that has a negative refractive index at microwave/radar wavelengths. The spacing between each lattice cell is 2.68 mm and the entire cube has sides of 2.5 cm (Image courtesy of Boeing Phantom works, Seattle)

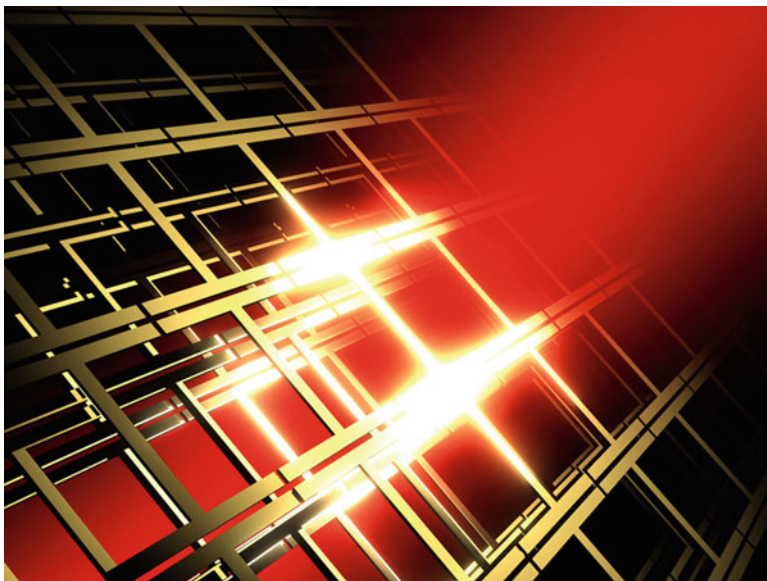


Fig. 5.12 Artist's impression of the I-shaped lattice configuration used to produce high (positive) refractive index metamaterials. In the prototype model, the I-shaped elements were $60\text{ }\mu\text{m}$ high and $60\text{ }\mu\text{m}$ wide (Image courtesy of Muhan Choi, Korea Advanced Institute of Science and Technology)

under a \$5.6 million DARPA (the Defense Advanced Research Projects Agency) contract, which started in 2001, to study and fabricate metamaterial structures and to investigate their potential application in the aerospace industry. The Boeing Cube was designed to work, that is, emulate a negative refractive index medium for incident radiation with wavelengths corresponding to several centimeters – i.e., the radar wavelength range.

The specific properties of a metamaterial are determined by the size, profile and spacing of its resonator elements. And, indeed, as we shall see, metamaterials need not always have a negative refractive index characteristic. By using a closely packed grid of I-shaped elements, for example, researchers in Korea announced in early 2011 that they had created a metamaterial with a record-breaking (positive) refractive index of $n=38.6$ – this is nearly ten times larger than the largest refractive index found in natural materials. The new metamaterial (Fig. 5.12) was described in the February 17, 2011 issue of *Nature*, and lead author Muhan Choi writes that the substrate was designed to operate at terahertz frequencies, corresponding to wavelengths of a few millimeters. Choi and colleagues also suggest that refractive indices of several hundred might eventually be achieved by reducing the spacing between the I-shaped elements and embedding them in a higher refractive index substrate. There are a whole range of possible uses for such high refractive index metamaterials, ranging from full body scanners (now being used at many airport security checkpoints), small-footprint invisibility cloaks (literally a black covering patch in this case), to medical physics and the diagnosis of various skin cancers.

The Harry Potter Project

Having worked through the theoretical details of the artificial magnet, Pendry realized that if the magnetic properties of a metamaterial containing tiny copper loops could be combined with the electrical properties of Marconi's radar-absorbing carbon compound then a way of manipulating electromagnetic radiation directly might be realized. With the right choice of composite materials the electromagnetic wave could in principle be guided and channeled to any desired location. Pendry and co-workers began working through the devilishly complicated details.

The story now shifts across the Atlantic Ocean. If the theory for invisibility cloaking came out of the UK, then it was in the United States that the impetus and money materialized to make them a reality. Following the publication of Pendry's work on magnetic metamaterials, along with Smith's and co-workers' early verification of the negative refractive index properties of metamaterial constructs, DARPA set up a conference to consider where this new technology might go. It was the year 2000 – the start of a new century, and the beginning of a new science.

Money was available, and on the table, for researchers to apply for, and interest was running high. Other conferences followed, but another twist to the story took place in 2004. This time DARPA had invited Pendry to give a review of the work being conducted at Imperial College, and to spice up his talk a little, Pendry decided to mention the so-called Harry Potter Project. It wasn't Harry's invisibility cloak, as such, that Pendry referred to in his talk but the secret, indeed, invisible entry point to Platform 9¾ at King's Cross Station – where access to the Hogwarts Express is gained. This part of the Harry Potter story is exactly the sort of technology that has military application – when is a door not a door? Why, when it is made to look like something else, of course. If a person can't see a door then they won't ask what lies behind it. Pendry recalls that he was offered large sums of money (almost on the spot) to carry on the research. The rabbit, so to speak, was out of the hat, and all of sudden talk of manufacturing invisibility cloaks became the rage.

Meanwhile, some 600-km north of Pendry's London workplace, Ulf Leonhardt was busily at work in his University of Saint Andrews office (Fig. 5.13). The year was 2005 and Leonhardt was not well pleased. The reason for his displeasure related to a research paper that kept being rejected for publication. The work had been in progress for many years, on and off, and related to how an invisibility cloaking device might be designed; laying out the mathematical techniques required.³ The paper had been rejected by such journals as *Nature*, *Science*, *Nature Materials* and *Physical Review Letters*, with the various referees making their decisions for reasons ranging from the work being too far fetched, to not offering enough new physics. The situation was frustrating, but one of the referees let slip an important piece of news. Leonhardt was not alone in his interests concerning invisibility cloaking.

³For those readers wishing to delve into the full-bodied mathematical details then we thoroughly recommend Leonhardt's recent book, written in collaboration with Thomas Philbin, *Geometry and Light – The Science of Invisibility* (Dover Publications, 2010).



Fig. 5.13 Ulf Leonhardt (b. 1965) whose research on conformal mapping made the design of invisibility cloaks practicable (Image courtesy of the University of St. Andrews)

The Imperial College group, lead by John Pendry, were also working on the topic and in fact were nearly ready to publish their work. When Pendry and co-workers submitted their paper to *Science* it received favorable reviews and was accepted for publication. At this stage the editors of *Science* realized that Leonhardt's previously rejected paper drew the same conclusions as the article just accepted from the Imperial College group. To their credit, the *Science* editors contacted Leonhardt and said they would like to publish his paper alongside that of Pendry's group. In May of 2006 the two papers appeared, side by side, in the press.

To the uninitiated reader it would not necessarily have been clear how important or indeed revolutionary the two papers actually were. Leonhardt's paper was the first of the two in the journal's pages and had the title *Optical Conformal Mapping*.

If the title left the reader unsure as to the paper's direction, then the first sentence of the abstract made matters clear – Leonhardt writes, “An invisibility device should guide light around an object as if nothing was there, regardless of where the light comes from.” The paper by Pendry, David Schurig and David Smith appeared immediately after Leonhardt's, and had the more understandable, albeit understated, title of *Controlling Electromagnetic Fields*. Again, the abstract to this paper made its topic clear, and the authors write, “Using the freedom that metamaterials provide, we show how electromagnetic fields can be redirected at will and propose a design strategy”. The abstract concludes, “Our work has relevance to exotic lens design and to the cloaking of objects from electromagnetic fields.”

With the theory now in place the next job was to make a demonstration model. Soon after Pendry's 2004 DARPA talk, his earlier college and collaborator Dave Smith (now at Duke University) decided to take-on the challenge. In order to make the fabrication process more manageable Smith and his team of researchers decided to work at microwave wavelengths rather than the much shorter and technically challenging wavelengths of visible light. In 2006 the first invisibility cloak (to be described shortly) was fabricated. The Duke cloak was made of a specially constructed metamaterial and successfully steered microwave radiation around a central cavity containing a copper cylinder. The properties of this device will be described in a moment, but in the meantime there are still a few more points to relate with respect to the early development of metamaterials.

Superlenses

We are back in the year 2000. After completing his work on the artificial magnet metamaterial, Pendry and co-workers realized that the negative refractive index properties of metamaterials could be used to make a superlens capable of revealing details on scales never before realized at optical wavelengths. The trick would be to use a metamaterial to funnel what is called the near-field radiation away from the object under investigation.

What the Imperial College group realized was that any source of electromagnetic radiation actually produces two types of electromagnetic wave – these are the so-called near-field and far-field waves. Importantly, the evanescent, near-field waves do not propagate into space like their far-field relatives; rather they die out very rapidly and are normally never seen. If these evanescent waves could be picked up, however, by the right kind of metamaterial magnifier, then they could be imaged. This would be of great benefit to microscopy since with respect to resolution the finest detail is always contained in the near-field waves – and it technically has no lower limit. In contrast, a standard lens can only work on far-field waves (in the geometric realm of ray diagrams) and the resolution for these waves can be no better than about half the wavelength of the light illuminating the object – corresponding to the so-called diffraction limit.

With a superlens, working on near-field waves, the detail that might be resolved is smaller than the wavelength of the illuminating light (Fig. 5.14). Such capabilities

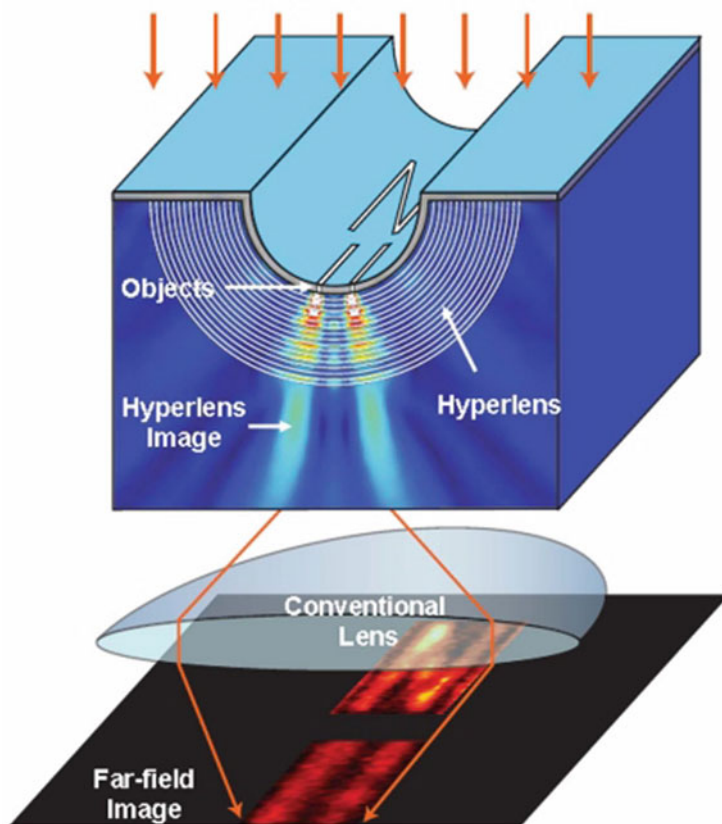


Fig. 5.14 Hyperlens design by Xiang Zhang and co-workers at the University of California, Berkeley. The hyperlens is made of alternating layers of silver and aluminum-oxide, which transport the high-resolution near-field (evanescent wave) image into the far-field domain, where it can be viewed with a conventional optical lens. The image at the *bottom* indicates that the hyperlens is able to resolve detail smaller than the wavelength of the illuminating ultraviolet light, beating, in fact, the diffraction limit of far-field optics (Image courtesy of Laurence-Berkeley Laboratory)

would, for example, revolutionize biology and plant sciences, enabling researchers to see live viruses infect living cells, or DNA strands replicating in real time. Detail on this minuscule scale has certainly been observed before, but only with devices such as scanning electron microscopes, which require that the specimen be specially coated and observed in a vacuum – conditions largely detrimental to the preservation of life.

The properties of a superlens working on far-field light waves were first described by Veselago in his 1968 ground-breaking paper. The crucial step, he realized, was to first construct an $n = -1$ metamaterial. If the lens substrate has a thickness d , then an object placed a distance $d/2$ away from one face will be imaged at a focal point a distance $d/2$ from the opposite face (Fig. 5.15). In contrast to the far-field case, Figure 5.16 shows how a superlens might capture and then transport for eventual viewing a near-field, evanescent wave. In this latter case the negative refractive index

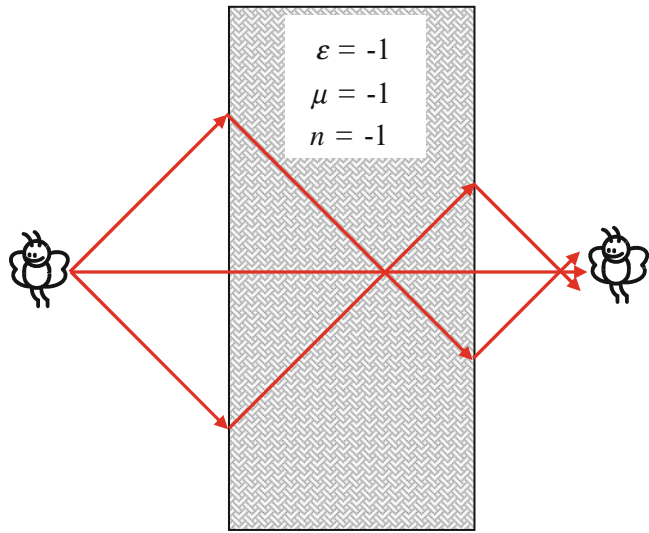


Fig. 5.15 The essential focusing operation of a superlens made of an $n=-1$ metamaterial slab

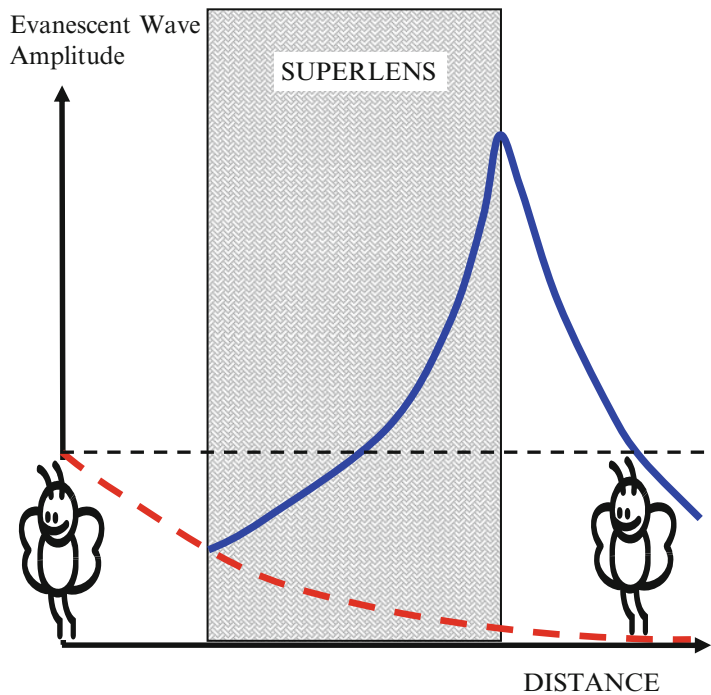


Fig. 5.16 Schematic illustration of the near-field, evanescent wave capturing characteristics of a negative refractive index metamaterial lens. In the absence of the superlens the evanescent wave would decay away rapidly (*dashed line*) and be unobservable with standard lenses

metamaterial substrate is positioned as close to the object to be imaged as possible, where it can pick up the otherwise rapidly decaying evanescent waves. Once captured, the metamaterial substrate can then channel radiation to a region where it can be viewed (see the upper left-hand corner of Fig. 5.3) by more traditional means.

It turns out that a superlens need not be a true negative refractive index metamaterial to work, but the minimum condition is that the electrical permittivity ϵ be negative. Indeed, a layer of silver just 40 nm thick was the active component of the very first optical wavelength superlenses to be made.

Again, with the theory in place, it was time for the experimental physicists and engineers to pick up the gauntlet. Researches at the University of Toronto in Canada led the charge and produced the first operational superlens in 2004. The experiment was conducted by Anthony Grbic (now at the University of Michigan) and George Eleftheriades (University of Toronto) and published in the journal *Physics Review Letters*. The metamaterial superlens transmission grid of the Toronto experiment operated at microwave wavelengths, and proved that sub-wavelength detail could be resolved. Superlenses were indeed possible, and once again the theoretical predictions were found to be right on the money.

Other research groups were hot on the heels of the Toronto team, and the newly established metamaterials laboratory at the University of California, Berkeley, headed up by Xiang Zhang, developed, in 2005, one of the first optical superlenses with a resolution many times better than that of the best available optical microscopes. A second research group situated half a world away, headed up by Richard Blaikie at the University of Canterbury in New Zealand, ostensibly employing the same techniques as the Berkeley group, also produced a superlens in 2004 to 2005. In 2007, however, the Berkeley researchers went one step further and fabricated a superlens (Fig. 5.14) that could resolve the 130 nm gap between two nanowires – a gap appreciably smaller than the wavelength of the illuminating ultraviolet light ($\lambda = 365\text{-nm}$). Detail that was previously invisible to our eyes, and even to state-of-the-art optical microscopes, was thus made visible and brought into the realm of the directly observable – an incredible advance in less than a decade of research.

The next step for optical wavelength superlens microscopes will be to couple them to a direct imaging system. In this respect Maxwell's classic fish-eye lens (recall Chap. 3) might potentially pave the way forward. Indeed, in the September 2009 issue of the *New Journal of Physics*, Ulf Leonhardt showed that if a Maxwellian fish-eye could be fabricated then a perfect image should be formed without the use of a negative refractive index substrate – a hitherto unexpected result. While Maxwell's fish-eye lens does not actually magnify, it can pick up, channel and bring to a perfect focus the troublesome evanescent waves that ordinary lenses miss. In practical terms, as pointed out by Leonhardt, the resolution of a Maxwellian fish-eye is limited to the size of the sub-wavelength structures within the lens itself and by the ability to accurately fabricate the required refractive index profile. Leonhardt suggested that molecular scale resolution (that is, scales as small as 10^{-10} m) might be achieved with a Maxwellian fish-eye lens in the not too distant future.

Leonhardt's 2009 paper on Maxwell's fish-eye lens caused some considerable controversy. A number of researchers, John Pendry among them, claimed that the

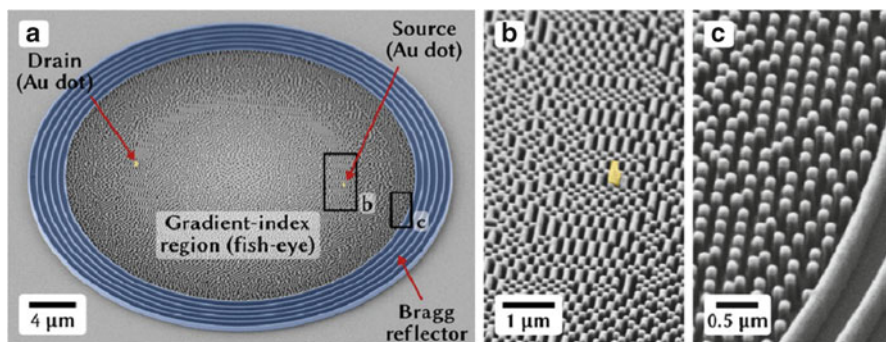


Fig. 5.17 The first Maxwell fish-eye lens. (a) Measuring just 30 μm across the refractive index of the lens is controlled by the careful placement of microscopic holes (b) and small silicon wires (c). The lens transfers a perfect image from the source (which can be located anywhere within the lens) to its symmetrically opposite point (the drain) relative to the center of the substrate (Image from arXiv:1007.2564v1)

analysis was flawed and that it would not be possible to fabricate a Maxwellian superlens – either in practice or in principle. A series of critical commentaries were published, with Leonhardt gamely responding with detailed rebuttals. The situation has still not been resolved to everyone’s satisfaction, but Leonhardt has collaborated with several research groups to actually fabricate Maxwell fish-eye lens substrates (Fig. 5.17) for direct experimental testing.

The first optical wavelength fish-eye device was constructed after Leonhardt teamed up with Lucas Gabrielli and Michael Lipson of the Nanophotonics Research Group at Cornell University, the results and preliminary testing of the miniature lens being published online at the arXiv webserver on July 15, 2010. On the very same day, a second research paper was added to the arXiv server concerning a second experiment in which Leonhardt, in collaboration with researchers from universities in Singapore and the Czech Republic, reported on the perfect imaging properties of a Maxwellian fish-eye lens fabricated to work at microwave wavelengths. In spite of these impressive experimental results, few of the original critics to Leonhardt’s ideas have been won over. Commenting to journalist Jon Cartwright (for physics-world.com – March 9, 2011), however, an undaunted Leonhardt observed that, “It will take time and more experiments, but I’m sure in the end even the most hard-nosed critics will be convinced... Seeing is believing.”

It is difficult to say when a revolution, scientific or otherwise, actually begins; it is like a giant beast, slow to awaken but voracious once fully aroused. Figure 5.3 indicates that within 5 years of the first negative refractive index metamaterial being demonstrated over 300 papers had referenced Veselago’s paper (after a 30 year slumber) – the subject of negative refractive index metamaterials might well be said to have fully awakened by this stage. Long live the revolution! With all this being said, however, we are still in the early days (read years to decades in practice) of the full establishment of this new technology. In the sections that follow we shall take a look at a few of the pioneering experiments, and marvel, as is only appropriate, at the magical devices that the new revolutionaries have already contrived.

The First Vanishing

The very first metamaterial invisibility cloaking device was manufactured in the Duke University lab of David Smith (Fig. 5.18). Announced in the October 19, 2006, issue of *Science Express*, the cloak in fact evolved from their earlier negative refractive index metamaterial experiments, and again, for ease of initial construction it was designed to operate at microwave wavelengths. The basic working principle of a cloaking device is illustrated in Fig. 5.19, with the circular metamaterial arrangement being used to channel electromagnetic radiation around a central cavity in which a solid object can be hidden.

The invisibility cloak built by Smith and co-workers is fabricated from a series of concentric circular substrates (Fig. 5.20). Etched upon the surfaces of these substrates are a series of split-ring resonators that produce the required radial variation in the negative refractive index, and which guide any incident microwave beam around the central cavity – thereby making the central region invisible to microwave detection. Each substrate is 1 cm high and the outer diameter is 13 cm – about the same diameter as the familiar music CD or DVD.

Moving towards slightly shorter operating wavelengths, Elena Semouchkina and co-workers at Michigan Technological University (MTU) introduced a remarkable infrared wavelength cloaking device in the summer of 2010. Writing in the June issue of the journal *Applied Physics Letters*, Semouchkina and colleagues describe the design of an invisibility cloak, made of glass,⁴ which can hide an embedded metal cylinder 1.5 cm in diameter (Fig. 5.21). The MTU design relies upon the emplacement of identical, nano-sized, nonmetallic chalcogenide glass cylinders

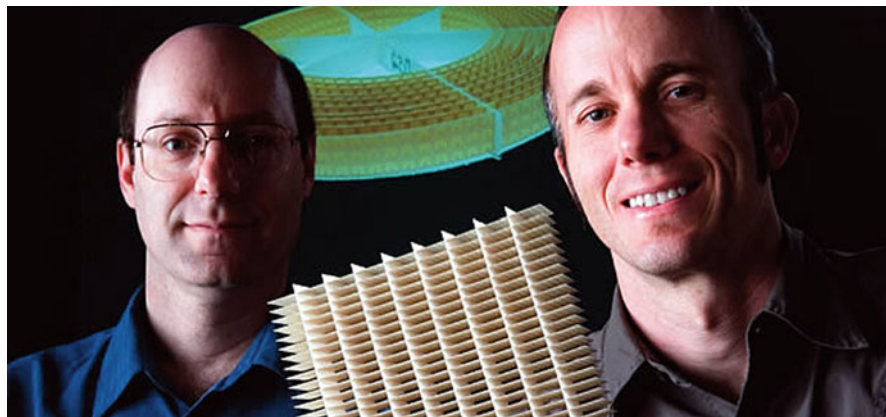


Fig. 5.18 David Smith (*left*) and David Schurig (*right*) of Duke University hold up part of a split ring resonator metamaterial substrate. The first microwave invisibility cloak is imaged in the background (Image courtesy of Duke University)

⁴There is a pleasing Charles Dodgson, Cheshire cat kind of irony to this design, in that an invisibility cloak is manufactured from what is quintessentially the most transparent of optical media.

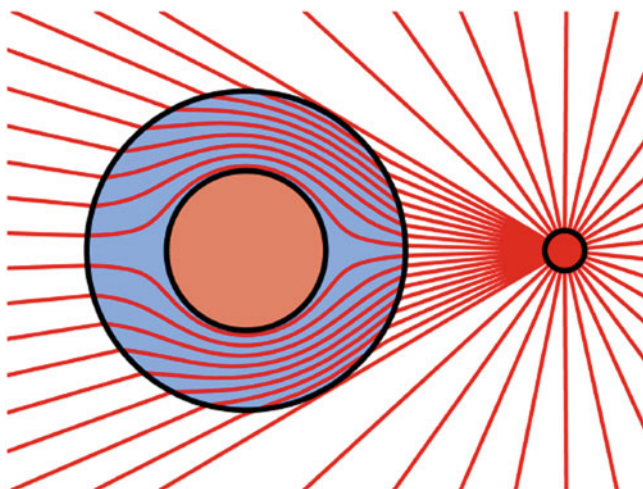


Fig. 5.19 Schematic ray diagram for the operation of an invisibility cloak. The electromagnetic radiation source is shown as the small circle to the *right*, and the cloaking device corresponds to the two concentric circles to the *left*. The key requirement of the cloak is that the incident radiation is channeled around its central core with no down-range observable effects – the radiation essentially flows around the cloaked region and returns to its original radial path (Image courtesy of D. Schurig)

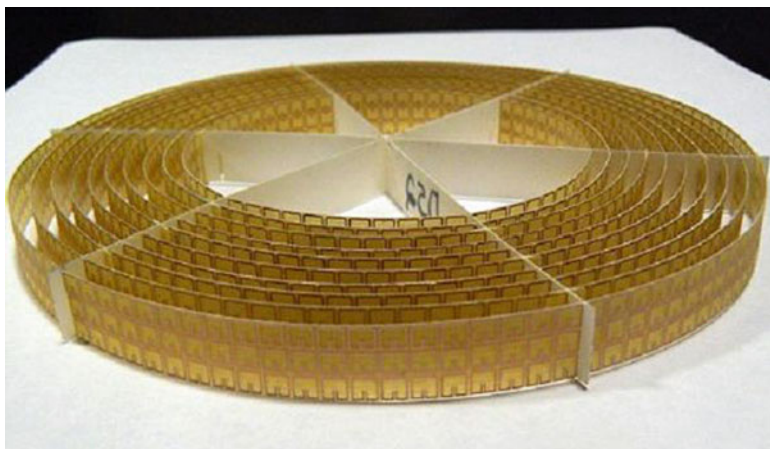


Fig. 5.20 The metamaterial microwave-cloaking device developed by the Duke University researchers. The cloaked inner region is 5 cm in diameter (Image courtesy of Duke University.)

within the supporting glass substrate. The embedded chalcogenide cylinders, which are essentially dielectrics with high refractive index characteristics, are arranged in a spoke-like, concentric circle pattern around the region to be cloaked. The cylinders are just 300 nm in diameter and 150 nm long. The MTU team has also developed

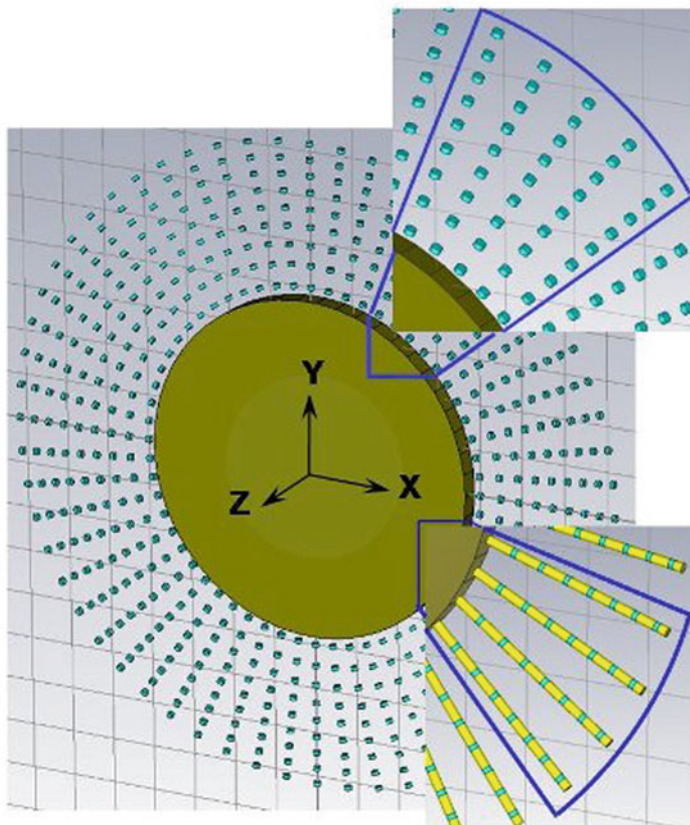


Fig. 5.21 Schematic of the MTU invisibility cloak made of glass. The *upper inset* indicates the spacing of the resonator arrays composed of chalcogenide glass, while the *lower inset* shows the radial spoke-like arrangement of glass resonators and fused silica glass (Image courtesy of Semouchkina and MTU)

a cloaking system made of ceramic resonators that operate at microwave wavelengths and with these they have successfully cloaked metal cylinders up to 7 cm in diameter and 10 cm high. The team's next stated goal is to move towards the construction of an optical wavelength cloaking system.

The Duke and MTU invisibility cloaks prove the capabilities of metamaterials and indeed, they also prove the concept of invisibility cloaking. The next step, indeed, the Holy Grail of present metamaterial research, is to produce a three-dimensional invisibility cloak operating at optical wavelengths. While in principle the physical design doesn't need to change when moving to visible light, the scale of the device, as well as its structural components, must. Indeed, the metamaterial substrates will have to take on dimensions smaller than that of the wavelength of visible light itself. This engineering step is not going to be an easy one, and indeed it will require a whole new manufacturing approach.

Nanomaterials and a Small Slice of Black

The nanometer is equivalent to one-billionth of a meter, or a millionth of a millimeter. Clearly, by human standards, this is small scale indeed; imagine slicing the following dash (which is about one millimeter in length) – into one million pieces. Compared to the atom the nanometer is still very large, but it is smaller by a factor of about 1,000 to the wavelength of visible light.

Nanotechnology is all about building substructures from the bottom up, atom by atom, molecule by molecule. The term for this kind of engineering was coined by pioneering researcher Eric Drexler in the 1980s, and since that time many national, industrial and university laboratories have been established to develop the techniques to manipulate nanoscale structures. These same laboratories are looking at innovative ways of manufacturing such diverse entities as new medicinal drugs, smaller computer chips, more compact electronic data storage devices, and even the construction of nanobots – the latter being in some cases the literal fabrication of minuscule mechanical robots.

In an ahead-of-its-time lecture presented by American physicist Richard Feynman at the California Institute of Technology (Caltech) on December 29, 1959, the idea of nanobots was introduced as follows, “I want to build a billion tiny factories [nanobots], models of each other, which are manufacturing simultaneously, drilling holes, stamping parts, and so on.” To this Feynman later added, “In the great future, we can arrange the atoms the way we want; the very *atoms*, all the way down.”

Incredibly, Feynman’s “deep future” is nearly upon us now, and while mechanical nanobots, as such, are still a distant dream it is now generally taken for granted that one day, not too far off, they will be deployed within human bodies to perform microsurgery and help monitor patient health. Indeed, the first molecular nanobots have already been programmed and manufactured. Writing in the June 20, 2010, online issue of the journal *Nature*, a team of researchers working at Harvard University announced the successful development of a self-assembling, three-dimensional nanostructure composed of single strands of DNA. The structure not only builds itself, it can also be pre-programmed to move and change its shape.

Perhaps surprisingly, the nanoscale alteration of glass has an ancient heritage, although the ancients did not understand the physics of the process – to them it was simply an art form. The manipulation in question is superbly illustrated by The Cup of Lycurgus⁵ (Fig. 5.22), one of the great treasures of the British Museum in London. Although the cup has been superbly crafted, it is the optical properties of the glass that is of interest here. Specifically, as illustrated in Fig. 5.22, when viewed via reflected light the cup appears to have a greenish tint to it. If illuminated from the inside, however, the cup takes on a reddish color.

The reason for this remarkable chromatic shift is a result of the ancient glass makers’ habit of adding fine metal grains to their glass-making recipes. In particular,

⁵Lycurgus (circa 800 B.C.), king of the Thracians.



Fig. 5.22 The Cup of Lycurgus. The cup was probably made in Rome about A.D. 290–325, and stands some 16.5 cm in height (Image courtesy of the British Museum)

when very small grains of silver and gold were added into the mix then so-called ruby glass was produced. This is the type of glass out of which the Cup of Lycurgus has been crafted, and it is an optical metamaterial. The light reflectance and transmission properties of ruby (or dichroic) glass are directly related to the size of the embedded metal grains – change the size of the grains and the color switch properties of the glass will change along with them. The remarkable result that we learn from this history lesson is that the manner in which light propagates through a layer of glass can be controlled by placing nanoscale wires within its interior.

The quest to make metamaterial substrates capable of manipulating light will strongly depend upon future developments in nanotechnology. This is entirely a size requirement condition, recall, since the substructures of the metamaterial need to be smaller than the wavelength of the incident radiation. Remarkably, however, nanowires and miniaturized split ring resonators (SRRs) have already been fabricated at several research labs (Fig. 5.23), and these developments will almost certainly lead the future charge in metamaterial research.

Realizing that there is an inevitable physical limit beyond which it will not be possible to fabricate SRRs, researchers at the Birck Nanotechnology Center at Purdue University have taken a lesson from the ancient Roman glass makers and outlined a possible method by which a non-magnetic metamaterial (i.e., one without SRRs)

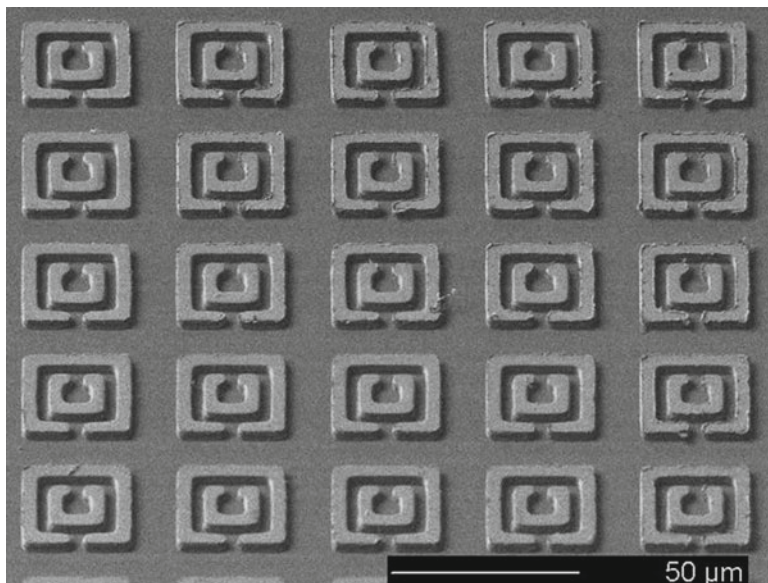


Fig. 5.23 An array of miniature split ring resonators produced through nanoengineering. The scale bar indicates that the individual structures are about 25 millionths of a meter in size (Image courtesy of Berkeley Labs)

might be fabricated to produce an optical wavelength invisibility cloak. Indeed, writing in the April 2007 issue of the journal *Nature Photonics*, Wenshan Cai and colleagues have described a way in which the Cup of Lycurgus might be transformed into an invisibility goblet, albeit at a very specific wavelength of light with a very specific polarization. The key to the process, they argue, is to carefully craft and insert nanowires into a cylindrical dielectric substrate. The group conducted a series of computer simulations to learn how such a cloak might work and found that by carefully adjusting the radial distribution of the nanowires an incident light ray could be smoothly directed around a virtual central cavity (Fig. 5.24). In essence the wires transform, as far as the light is concerned, a solid glass cylinder into a cylindrical glass shell. Cai and co-workers conclude that their method “paves the way for the realization of cloaking devices at optical frequencies,” but to date no such cloak has been fabricated, and the theory currently runs ahead of the art.

Total Absorption

The metamaterial approach to creating stealth technology took a major step forward in the summer of 2010 with the announcement, by (another) Purdue University researcher named Evgenii Narimanov, that a new light-absorbing surface had been engineered. Building upon the ideas outlined by John Pendry, the

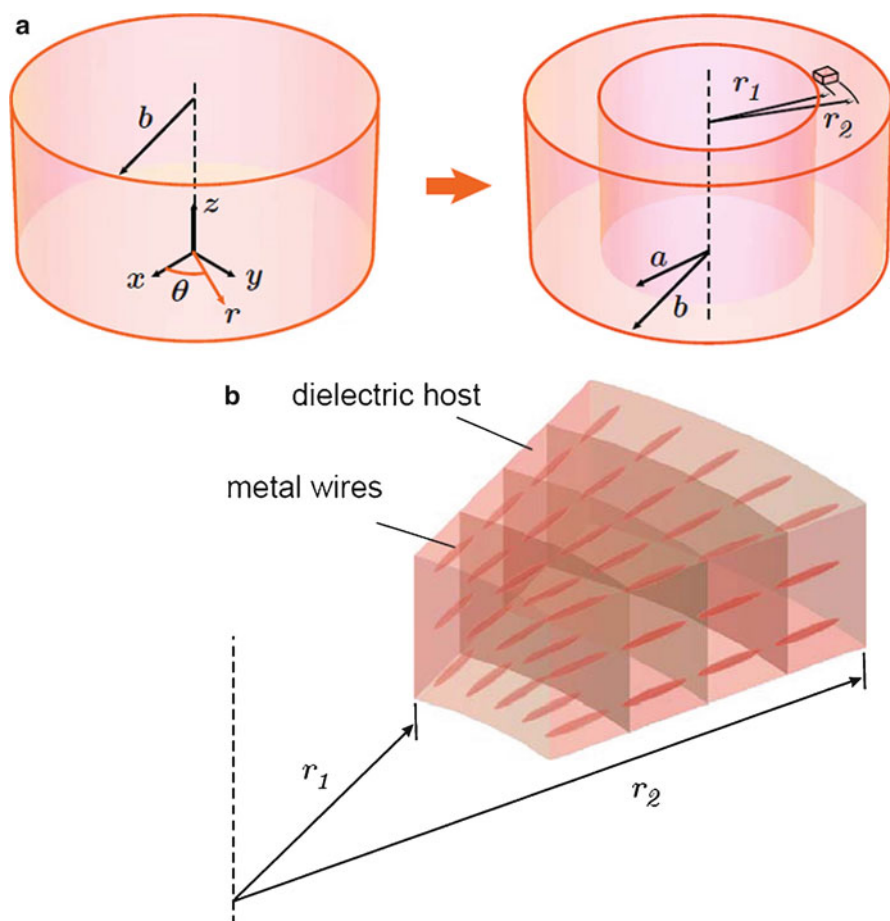


Fig. 5.24 Outline of the Purdue optical cloak. (a) Schematic view of the transformation being invoked – a solid dielectric glass cylinder of radius b is transformed so as to appear to a light ray as a cylindrical glass shell with an inner core radius a . An object can in principle be hidden from view within the inner core region. (b) An enlarged slice of the cloak, showing the radial distribution of the all-important metal nanowires (Image from arXiv:physics/061124v1)

new metamaterial is made of thin silver wires (Fig. 5.25) just 35 nm in diameter. Narimaov and co-workers studied the properties of their newly constructed metamaterial at infrared wavelengths and found that for a smooth surfaced sample some 20% of the incident radiation was reflected. This, of course, is far removed from the stealth technology standard. However, in a nice example of where perfection is far from ideal, the researchers found that less than 1% of the incident light was reflected if the surface was roughened. Under these circumstances the metamaterial truly was a good absorber – at least at infrared wavelengths. The next step is to expand the wavelength region over which the metamaterial works. Narimanov

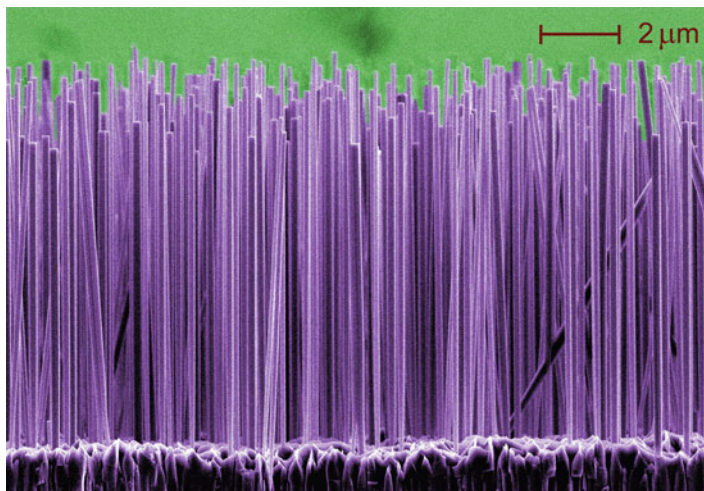


Fig. 5.25 Nanowires can literally be grown on a substrate through a process of vapor to solid deposition. This image shows nanowires of a gallium nitride alloy grown at the National Institute of Standards and Technology in Gaithersburg, Maryland. These particular fibers can emit ultraviolet light and are part of a project to build nanoscale lasers. The scale bar in the *upper right* corresponds to two millionths of a meter (Image courtesy of NIST)

is confident that a truly universal, all wavelength absorbing metamaterial will eventually be engineered, but the day of its actual realization is located somewhere in the distant future.

An alternative, non-nanowire design for an all-wavelength absorber has also been described by Qiang Cheng (Southeast University, Nanjing in China) and co-workers in the June 2010 issue of the *New Journal of Physics*. Their device is an omnidirectional absorber composed of 60 concentric rings of metamaterial substrates (Fig. 5.26).

The omnidirectional electromagnetic absorber has been fabricated so that any incident light falling upon its outer rim is channeled via a spiral path towards its center, where it is eventually fully absorbed—in essence it is an electromagnetic black hole. Unlike an astrophysical black hole, however, the metamaterial absorber has no central singularity, or gravitational effect, but it does mimic some of the properties that will appear near a black hole’s event horizon – that trap-door boundary beyond which there is no escape (recall Chap. 1 and see Appendix 1).

In its present form the absorber constructed by Cheng and collaborators works at microwave wavelengths, and the energy carried by these waves is deposited within its central core. The result of this heating means that the absorber actually acts more like an electromagnetic blackbody radiator rather than a space warping black hole. The most immediate applications of this new metamaterial device are likely to be found in those areas where thermal emission and/or electromagnetic wave harvesting is of interest – such as in the development of more efficient solar panels for the generation of electricity.

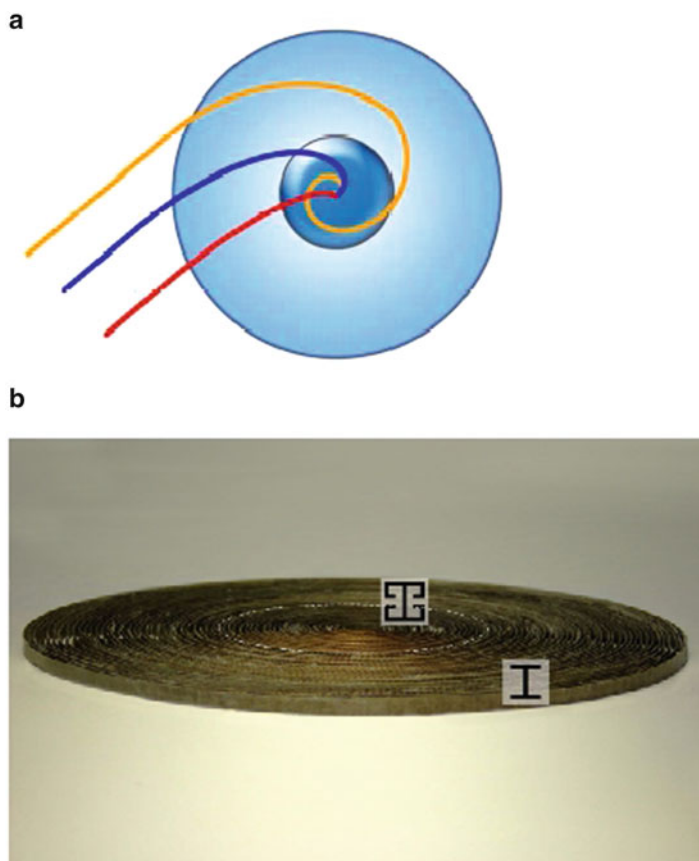


Fig. 5.26 The omnidirectional electromagnetic absorber analog of a black hole. The absorber **(b)** builds upon the design of standard cloaking devices, but in this application all the light is channeled towards the central region **(a)**. The *inset* images on **(b)** show the metamaterial unit cell designs (Image courtesy of Q. Cheng and NJP)

Carpet Cloaking

The essential idea behind the carpet cloak, introduced in Chap. 3, is to make an uneven reflective surface, covered in bumps and wrinkles, look like a flat, or at least benign, surface to an external observer. The objects to be hidden in this case are literally swept under an optical carpet.

The basic design requirements of a carpet cloak are illustrated in Fig. 5.27, where it can be seen that there are essentially two active components. The object to be hidden is placed under a smooth hump in a reflective, mirror-like base plate, and the distortions caused by the hump in the mirror are counteracted by an overlying meta-material substrate. To make an invisibility cloak work properly the refractive index

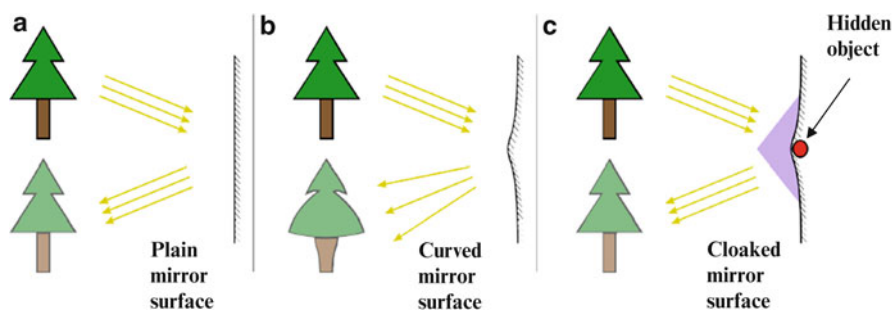


Fig. 5.27 Illustration of the carpet cloaking concept. (a) Reflection from the surface of a plain mirror shows no distortion in the image of the tree. (b) Reflection from a mirror with a hump produces a distortion in the reflected tree image. (c) Correcting the reflection for the mirror hump with an invisibility device – now a solid object (*circle*) can be hidden from view without the tree image being distorted (Image adapted from L. Gabrielli et al. 2009: <http://arxiv.org/abs/0904.3508>)

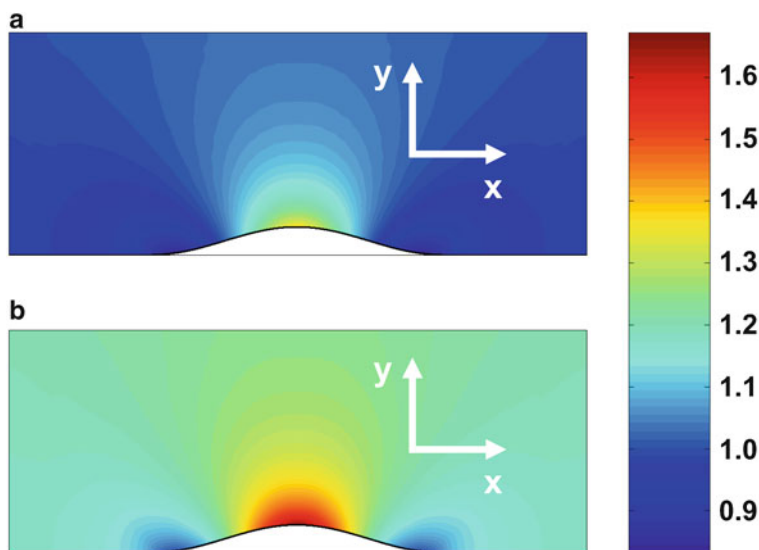


Fig. 5.28 False color images of the refractive index variation within carpet cloaking structures operating at optical wavelengths. (a) Variation with a background refractive index of $n=1$. (b) Variation with a background refractive index $n=1.2$ (Image courtesy of Karlsruhe Institute of Technology)

of the metamaterial layer is adjusted so that any incident light beams behave as if they have been reflected from the surface of a perfectly flat mirror.

Two examples of how the refractive index must vary inside a carpet cloak metamaterial are shown in Fig. 5.28 – Jad Halimeh and co-workers at the Karlsruhe Institute of Technology in Germany produced these particular images. Writing in the October 26, 2009, issue of the journal *Optics Express*, Halimeh et al. used a specially

written computer program to investigate the properties of an ideal carpet cloak. To compensate for the fact that a light ray reflected from the top of the bump travels a shorter path length than that of a light ray if no bump were present, the refractive index close to the bump must be higher, so that light travels more slowly there. In this way Fermat's least travel time (recall Chap. 2) condition is conserved. At the surface of the carpet and in regions well away from the bump the refractive index takes on a constant so-called reference index. If the reference index is larger than $n=1.2$ (Fig. 5.28 – image b) then at no point in the interior of the cloak need the refractive index be smaller than $n=1$, that is, smaller than the refractive index of air.

One of the first carpet cloaks to be fabricated was that by the Duke University research group in 2008. Working under the guidance of David Smith, the results from this groundbreaking project were published in the January 16, 2009, issue of the journal *Science*. The cloak was designed to operate at microwave wavelengths and was manufactured from metamaterial substrates supporting an extended array of H-shaped copper elements.

In this particular application it was not necessary for the metamaterial to produce a negative refractive index, as required in the earlier Duke invisibility cloak, but it did need to generate a smooth positional variation in the (positive) refractive index. Indeed, Smith and co-workers had to carefully calculate the size and separation of the metamaterial substrates, as well as the dimensions of the various etched-copper SRR elements. The experimental array was fabricated to cloak a smooth, dome-shaped region some 12 mm high and 125 mm long situated under a reflective copper ground plane mirror (Fig. 5.29). The entire cloak is 106 mm wide, 500 mm long and 10 mm high. A total of 10,000 individual 2×2 mm cells were fabricated to achieve the required refractive index variation for incident microwaves in the wavelength range from 19 to 23 mm. In this manner, when the angle of incidence was 40° , the microwaves were correctly directed around the bump to be reflected as if the ground plane were a flat surface.

In order to construct their microwave carpet the Duke University researchers used the theory of transformation optics (to be described below) to map out the required variation of the refractive index in the region of the cloaked bump. In this manner, as graduate student Ruopeng Liu, the lead author of the January 16 research paper, notes, “One imagines warping space so as to control the trajectories of light in a desired manner.” For the Duke cloak, the background refractive index was chosen to be $n=1.331$, but in the region of the bump it varies from $n=1.08$ to 1.67 .

At the same time that the Duke University group was developing their microwave carpet cloak, several other research teams were building and testing other devices which operated at near-optical wavelengths. Key to the success in fabricating these latter cloaks was the development and perfection of nano- and micro-scale engineering techniques.

The first optical wavelength carpet cloak was developed at the National Science Foundation Nanoscale Science and Engineering Center at the University of California, Berkeley. Indeed, working under the guidance of Center Director Xiang Zhang (Fig. 5.30), two research groups set about developing two very different fabrication programs. One group nanoengineered a wafer substrate made of alternating



Fig. 5.29 The microwave metamaterial carpet cloak built by researchers at Duke University in 2008. The cloaked region is under the central copper bump – David Smith is pictured in the background (Picture courtesy of Duke University)

layers of silver and magnesium fluoride cut into a fishnet array (Fig. 5.31). The second group engineered a substrate made of silver nanowires (recall Fig. 5.25).

The first team to publish from Zhang’s lab was that lead by graduate student Jason Valentine. This work appeared in the August 13, 2008, issue of the journal *Nature*, and described the fishnet carpet cloaking system. The fishnet array was fabricated from 21 layers of material alternating between silver and a non-conducting layer of magnesium fluoride. The characteristic squares of the fishnet were then cut into the layered substrate so that each pair of conducting and non-conducting layers could form current loops capable of interacting with the magnetic field of the incoming light. Valentine comments that the fishnet design “is the first bulk material

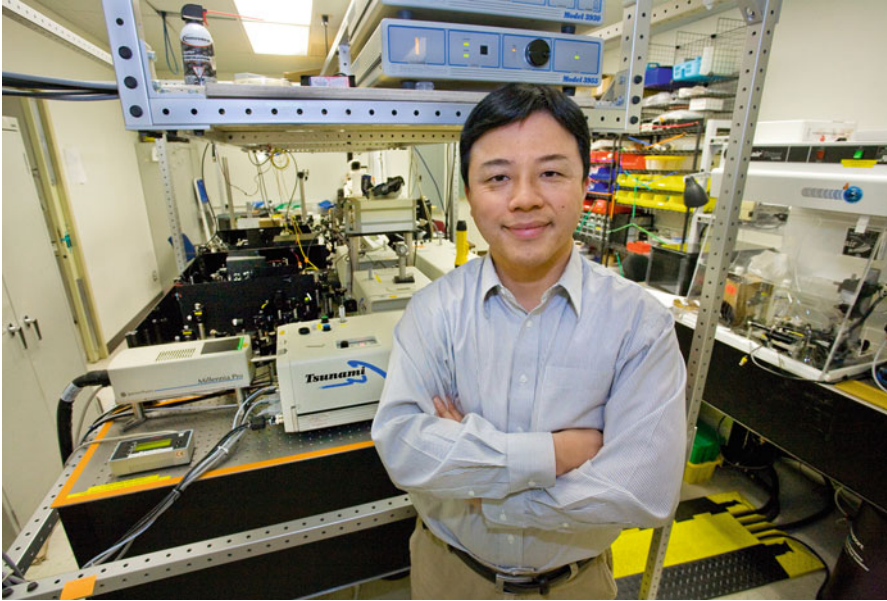


Fig. 5.30 Xiang Zhang, director of the NSF Nanoscale Science and Engineering Center at the University of California, Berkeley (Image courtesy of Berkeley Labs)

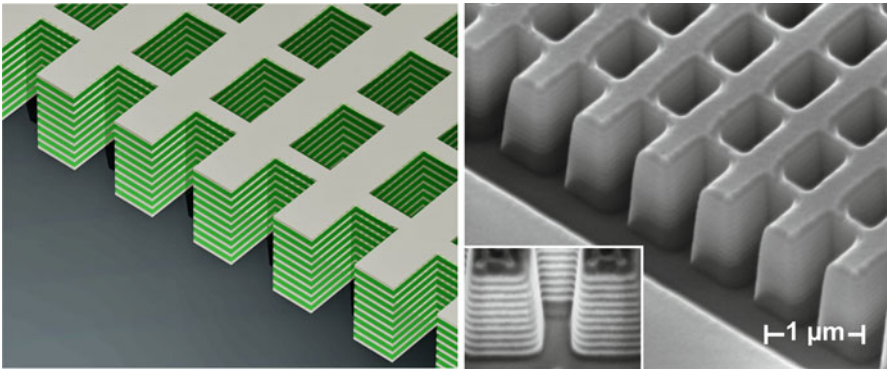


Fig. 5.31 Fishnet metamaterial developed at Berkeley. The design ideal is shown to the *left*, and the nanoscale fabrication is revealed in a scanning electron microscope image on the *right*. The scale bar (*lower right*) corresponds to a spacing of one-millionth of a meter (Image courtesy of University of California, Berkeley)

that can be described as having optical magnetism, so both the electrical and magnetic fields in a light wave move backward in the material” – this latter effect, as described earlier, being exactly the characteristics expected of a negative refractive index medium.

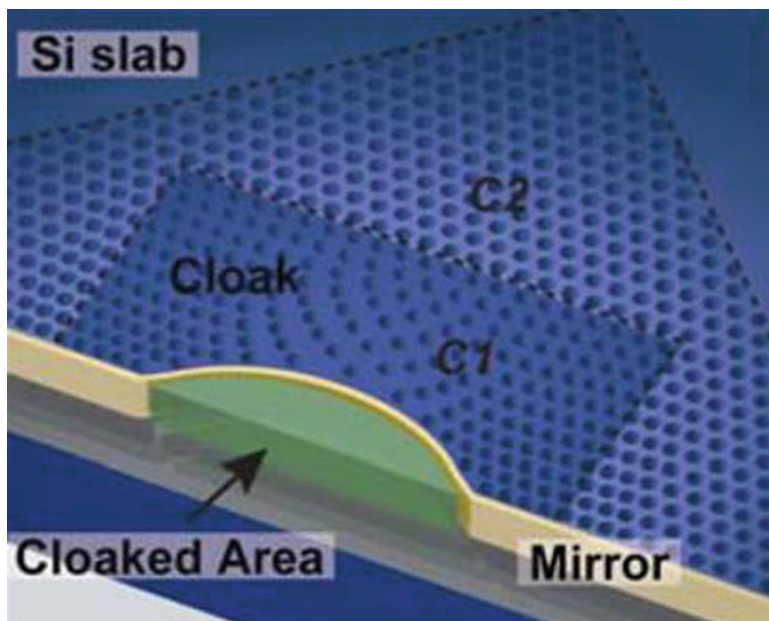


Fig. 5.32 Schematic of the Berkeley dielectric carpet cloak. The variable hole spacing in the region labeled C1 produces the required refractive index gradient to mask the cloaked region (shown in *green*). The reflective gold layer is shown in *yellow*. The regular spacing of the holes in region C2 provides a constant refractive index background. The light ray enters and exits the cloak at a constant angle perpendicular to the triangular edges. (Image courtesy of Berkeley Lab)

The second Berkeley team published their results, just 2 days after Valentine et al., on 15 August 2008 in the journal *Science*. Rather than using a fishnet matrix, the second group concentrated on manipulating the electric field of the incoming light ray. Their metamaterial substrate was comprised of equally spaced, parallel rows of silver nanowires that had been grown on an aluminum oxide substrate. Remarkably, the entire metamaterial structure is about ten times thinner than ordinary writing paper (although this is still ten times larger than the wavelength of light). While not actual functioning as a cloaking device the nanowire substrate was shown to have the required negative refractive index characteristic to produce one if so desired.

Not quite 1 year on from their first publications, Zhang's group at Berkeley published the details of yet another design for a carpet cloak. Writing in the July 19, 2009, issue of *Nature*, lead author Jason Valentine explained that the new cloak was fabricated from a triangular shaped "silicon-on-insulator (SOI) wafer consisting of a 250 nm thick silicon layer separated by a 3 μm thick silicon oxide slab from a Si [silicon] wafer substrate." A 100 nm layer of gold was applied to one side of the triangular carpet in order to act as the reflecting surface. The resultant, minuscule cloak is shown in Fig. 5.32, with the concealment region, below the gold reflecting surface, having a height and width of 3.8 μm and 400 nm, respectively. The refractive index of the substrate is determined by positioning a precise pattern of holes in the cloaks silicon

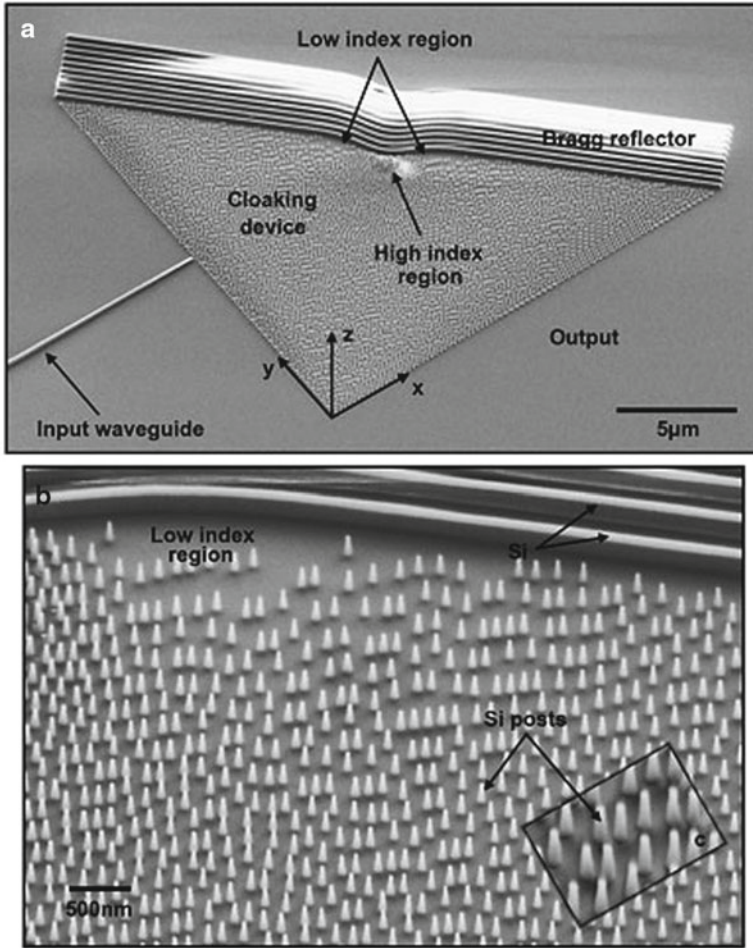


Fig. 5.33 Scanning electron microscope image of the Cornell carpet cloak. The upper image shows the basic structure of the device and the small 5- μm -wide cavity region to be cloaked. The lower image shows the “silicon posts” that have been etched onto a silicon-on-insulator wafer. Each post is 50 nm in diameter (Image courtesy of the Nanophotonic Group, Cornell University)

layer. Each hole is 110 nm across and is produced with a focused ion beam. The cloak described in the *Nature* paper was effective at hiding a carpet bump in the wavelength region from 1,400 to 1,800 nm, but Valentine and co-authors note that smaller wavelengths could be accommodated by reducing the diameter and spacing of the holes.

The research paper describing the Berkeley carpet cloak closed with a *note added in proof*, indicating that after their article had been accepted for publication the authors became aware of another device that had been fabricated by the Nanophotonic Group at Cornell University. This latter research by Lucas Gabrielli and co-workers was published online in *Nature Photonics* on July 20, 2009, where the construction of a small, triangular carpet cloak using silicon nanowire-like posts to direct near-infrared incident light around a cavity (Fig. 5.33) is described.

The carpet cloaking designs and new nano-metamaterials developed at the Berkeley and Cornell labs afford a significant operational improvement over the first split ring resonator cloaking devices. The key development is that the nanomaterials do not rely upon a specific resonance feature in order to operate (that is, to activate their negative refractive index characteristics). Accordingly they are able to function over a much broader wavelength range, and they are subject to much lower energy losses.

The properties of fishnet-style metamaterials have been further improved, however, by researchers at Purdue University through the introduction of a fluorescent dye (Rhodamine 800) layer sandwiched between the fishnet's silver and alumina substrates. Announcing their results in the August 5, 2010, issue of the journal *Nature*, Vladimir Shalaev and co-workers explain that the dye is "pumped" with a 690-nm wavelength light pulse. This pumping stimulates the dye to emit light at a wavelength of 730 nm for several nanoseconds. If in this brief emission time a second pulse of light, now with a wavelength of 730 nm, is fired at the dye an amplified probe pulse will result.

Since the dye has to be stimulated externally, Shalaev and colleagues call their new substrate an active negative-index metamaterial (which in our overly acronymic world they reduce to *active NIM*). The amplification afforded by active NIMs will greatly improve, Shalaev and team argue, the capabilities of superlenses and cloaking devices by reducing substrate energy losses to miniscule levels. Indeed, the Purdue researchers report a reduction in energy absorption losses by a factor of about 1,000,000 when the dye is activated compared to those times when it isn't. The Purdue team is currently looking at ways in which to replace the fluorescent dye with a semi-conductor material that can be "pumped" electrically.

In spite of the incredible advances in metamaterial fabrication there is still a very long way to go before a fully functional Harry Potter-style, flexible-cloth invisibility cloak, or even a carpet cloak that might hide a sugar lump, are going to be constructed. Not only must the fabrication scale be greatly increased, but so too must the wavelength range over which a cloak operates be expanded.

The carpet cloaking devices produced prior to late 2009 were only functional at one specific angle of incidence. If the incidence angle changed then the cloak no longer functioned correctly. In the hyperactive world of nanotechnology and metamaterial research, however, this viewing limitation was soon bested. Indeed, writing in the April 16, 2010, issue of *Science*, Tolga Ergin and co-workers at the Karlsruhe Institute of Technology in Germany described the successful fabrication of the first fully three-dimensional carpet cloak. The new style cloak (Fig. 5.34) used a grid of polymer rods, each just a few hundred nanometers in diameter, staked in layers that alternately lie across each other in a woodpile fashion – indeed, the device is called a woodpile photonic crystal. The refractive index of the various layers within the cloak is determined by a variable filling factor, and this was manipulated by a process known as direct laser writing. The refractive index was adjusted, between a background index of $n=1$ (corresponding to air) and a maximum of $n=1.5$ (corresponding to the bulk polymer). The entire structure was engineered to conceal a cavity one-millionth of a meter high and 13 millionths of a meter wide in a gold

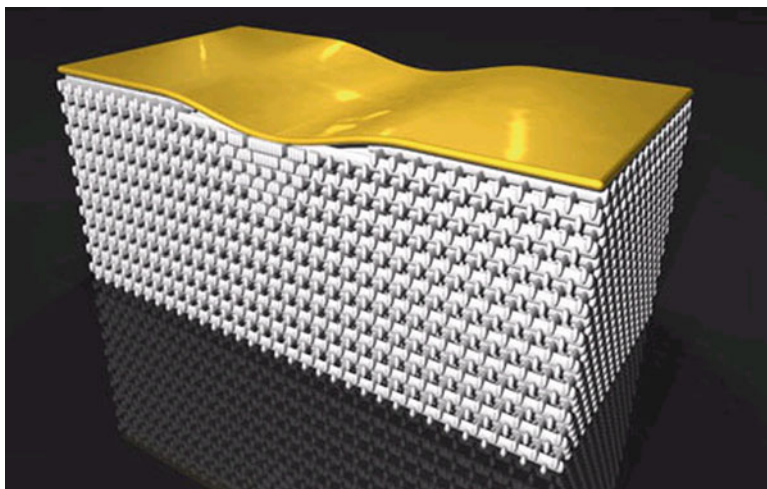


Fig. 5.34 Computer-generated view of the 3-D carpet cloak fabricated at the Karlsruhe Institute of Technology. The filling factor between the substrate rods increases near to the cavity to be cloaked – the mapping of the filling factor and the associated refractive index variation were calculated by transformation optics theory. The entire cloak has dimensions of $26 \times 10 \times 10 \mu\text{m}$ (Image courtesy of the Karlsruhe Institute of Technology)

ground plain reflector. The efficacy of the cloak was demonstrated over a wavelength range from 1.4 to $2.7 \mu\text{m}$ (still outside of the human eye range, but edging closer), and over a 30° angle range relative to vertical.

While the Karlsruhe group was first to go into print with their 3-D carpet design, the day before their article appeared in *Science* another research group uploaded a preprint of their work on a 3-D carpet cloak to the arXiv online data server. This team, headed up by Cheng Sun (Northwestern University, Evanston, Illinois), manufactured a triangular-prism shaped cloak that operates at microwave wavelengths. The cloak has a base length of 8.5 mm and is 4.2 mm high, and conceals a bump 0.5 mm high and 2 mm long. The NWU cloak was fabricated through a technique known as projection microstereolithography, and in this manner 200 layers, each $20 \mu\text{m}$ thick, were produced in a photocurable polymer resin. The interior of the cloak is cut through by an array of small holes variously spaced so as to produce the correct refractive index variation.

Less than 2 months after the paper by the Karlsruhe group appeared in print yet another breakthrough paper appeared. Writing in the June 1, 2010, issue of *Nature Communications*, Hui Feng Ma and Tie Jun Cui (Southeast University, Nanjing, China) described the design and fabrication of a broadband, three-dimensional ground plane cloaking device. Working at microwave frequencies this 12.5 cm diameter cylindrical cloak works from any viewing angle and cloaks a cone-shaped region 1.3 cm high and 12.5 cm in base diameter. The actual cloak is composed of multiple stacked 1-mm thick dielectric disks (Fig. 5.35c). Holes of carefully machined size and surface density are cut in to each of the disks (Fig. 5.35a, b), and these act to control the refractive index variations from $n=1$ to $n=1.6$ (Fig. 5.35d).

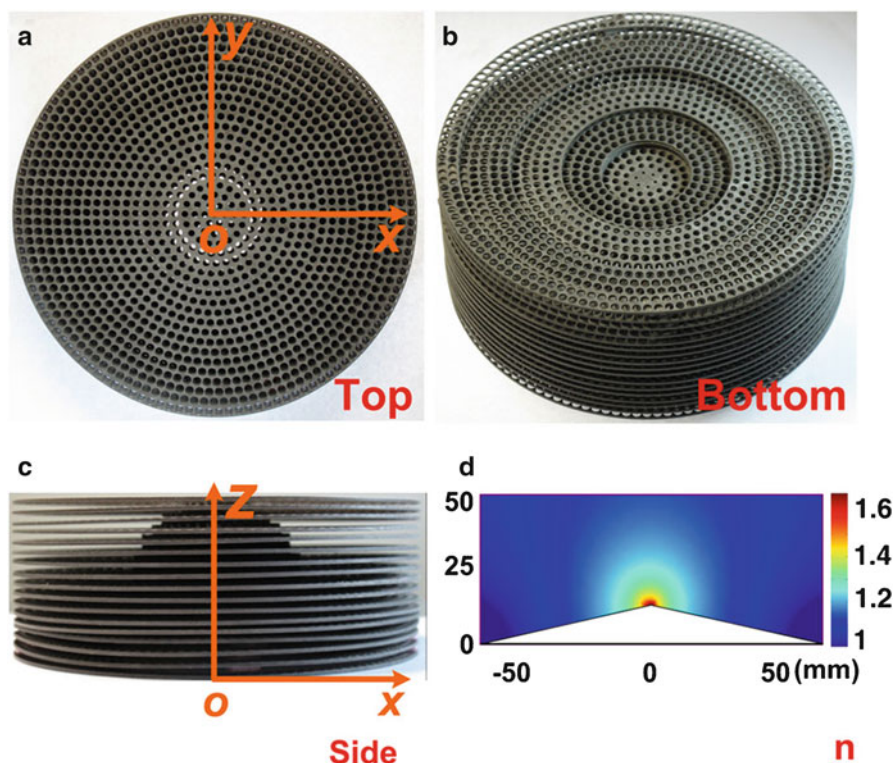


Fig. 5.35 Top (a), bottom (b) and side (c) views of the broadband, three-dimensional microwave ground plane cloaking device. (d) The refractive index profile for the cloak (Image courtesy of *Nature Communications*)

The cloak developed by Ma and Cul was found to be effective over the frequency range from 8 to 12 GHz, covering the so-called X-band region that is used by many modern radar installations. Ma and Cul conclude their paper with the hopeful, but yet to be fully realized, comment that, “owing to the easy fabrication [methods], the fully 3D cloak can be directly extended to the optical frequency.” Indeed, the cloak developed at Cornell University (Fig. 5.33) is a step in the direction indicated by Ma and Cul, but such devices are still far removed from being able to hide even millimeter-sized objects.

The metamaterials used in the carpet cloaks discussed so far have all been designed to remove a bump from a surface reflector, making it look like a flat mirror; there is no reason, however, to stop them from being designed to introduce artificial features. At least one research group has looked at the idea of constructing a remote, optical illusion cloaking device. Indeed, a hybrid of the Trojan horse design has been studied by Chi Ting Chan and co-workers at the Hong Kong University of Science and Technology. Writing in the June 22, 2009, issue of *Physics Review Letters*, Chan and collaborators describe the design of a cloak that can not only

make an object look like something else, it can also do it remotely, not necessarily being in contact with or covering the object to be disguised. Chan et al. write, “The principle behind this illusion is not light bending, but rather the exact cancellation and restoration of the optical path of light.” Perhaps a more prosaic way of saying this is that the cloak in this case steals the shadow cast by an object and replaces it with the shadow of something else – an apple can literally be made to look like a banana in shadow world. The device described by the Hong Kong University group is purely theoretical at this time, but in the way that science fact is typically just science fiction plus a good dose of time and research so, too, cloaking reality appears to be just a matter of metamaterial cloaking theory plus laboratory time – someone, somewhere will take up the challenge and turn the theory into practice.

And, finally – it had to happen. When is a door not a door? Why, when you can’t see it, of course. Just 6 years on from the time that John Pendry described the Imperial College London “Harry Potter Platform 9 $\frac{3}{4}$ Project,” the first invisible corridor was constructed. The experimental work was described by Li Chao (Chinese Academy of Science) and co-workers in the December 3, 2010, issue of *Physical Review Letters*. The demonstration model was fabricated on a so-called transmission line substrate made of capacitors and inductors. The network design, however, effectively formed a corridor between two perfect electrical conductors (PECs), but added an inner lining of negative-refractive index metamaterial to one of the walls. In this fashion an invisible gateway was produced in the sense that (at certain wavelengths – radio waves in this case) there was no transmission through the corridor. In other words an observer, even if they were looking right at the corridor (with radio-sensitive eyes), would not be able to see through it. Indeed, the corridor would look just like the surrounding walls. Not letting their eyes deceive them, however, the observer could nonetheless physically walk through the corridor (Fig. 5.36).

Nothing Is Perfect

The best way to catch a poacher, so the old saying goes, is to employ a poacher as your game keeper. The same adage applies to the detection of invisibility cloaks, and Baile Zhang (MIT) is the game keeper of choice. Having gained his doctorate from MIT in 2009 by writing on invisibility cloaks, Zhang turned his thoughts to ways in which even a hypothetically perfect invisibility cloak might be detected. While much of the initial funding for cloaking research has (obviously) concentrated on finding ways of hiding objects, funds have also been made available to look at ways that an invisibility cloak might be revealed – Zhang’s work, for example, is partly funded by the U.S. Air Force and the Office of Naval Research.

To date, two methods have been proposed, both by Zhang and co-workers, but neither approach is easily realized in practice. The first method by which a perfect ground cloak might be detected is based upon looking for lateral image shifts, while the other approach depends upon spraying a target region with high velocity electrons.

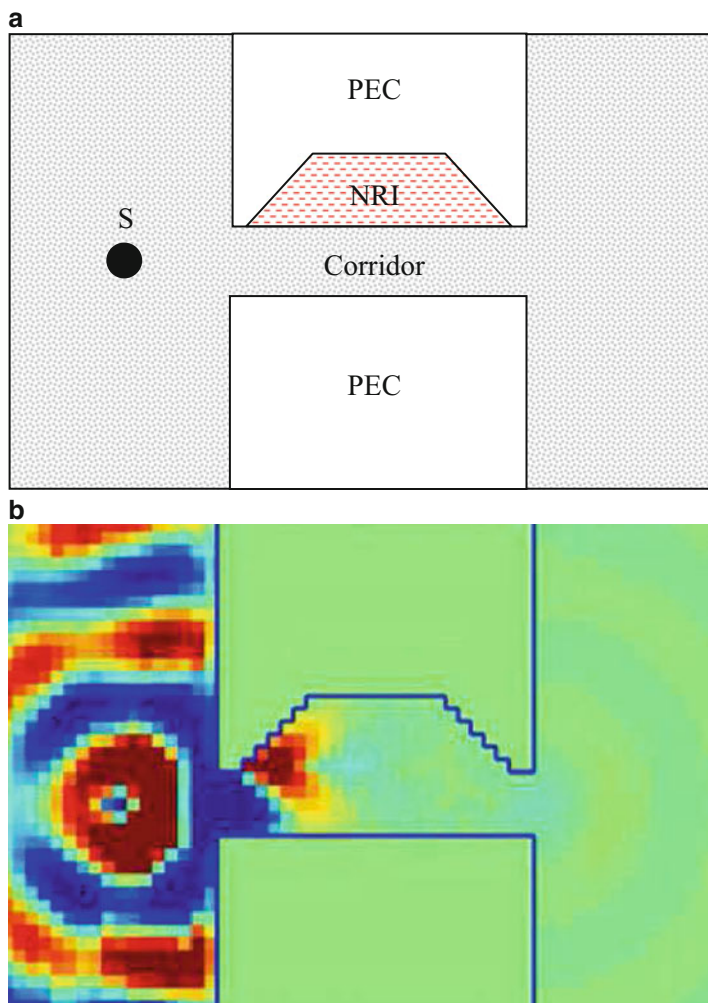


Fig. 5.36 (a) Schematic outline of the electromagnetic invisible gateway. Selective wavelengths of electromagnetic radiation from the source S will not propagate between the perfect electrical conductor (PEC) corridor because of the negative refractive index (NRI) lining. (b) Experimental realization of the invisible gateway – radiation from the source at *left* does not propagate through the corridor

Writing in the December 17, 2009, issue of *Physics Review Letters*, Zhang and co-author Bae-Ian Wu (also at MIT) explain how fast-moving electrons can be used to reveal even a perfect, all wavelength, three-dimensional cloaking device – at least in theory. The method of detection plays upon the natural behavior of charged particles to radiate electromagnetic radiation when their velocity changes. In this manner, when an electron encounters the metamaterial shell of the invisibility cloak, although it will eventually pass right through it, its velocity will nonetheless be changed, and this will produce a spray of so-called Bremsstrahlung, or breaking

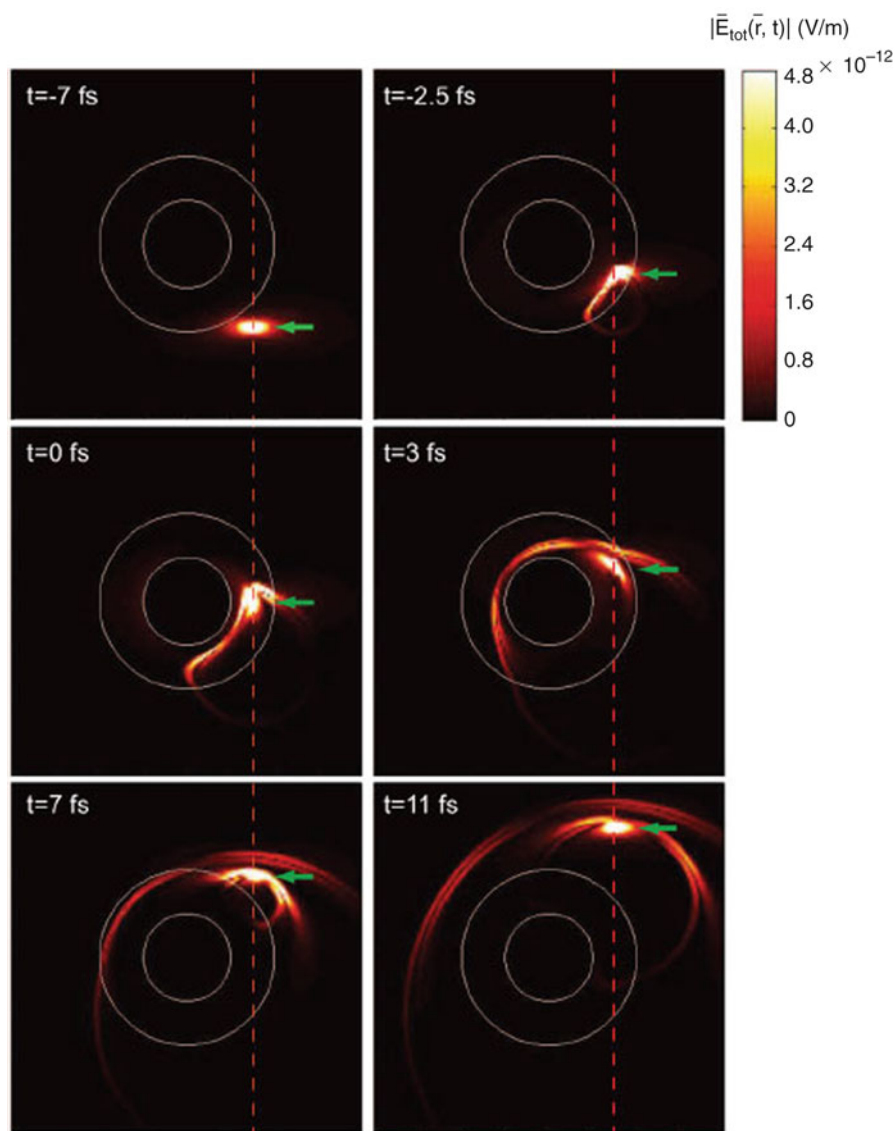


Fig. 5.37 Telltale radiation (shown in *orange*) emitted by an electron as it moves along a trajectory (*dashed line*) taking it through a spherical invisibility cloak. Some of the radiation is lost from view within the cloak, but importantly some also leaves the cloaking region. The location of the electron as a function of time (shown in femto seconds) is indicated by the short *green arrow* (Image courtesy of MIT)

radiation (Fig. 5.37). This characteristic flash of radiation will be visible to an external observer and accordingly the cloaking device will have been betrayed. There is no hiding from this detection method, and in principle any electromagnetic cloak is detectable.

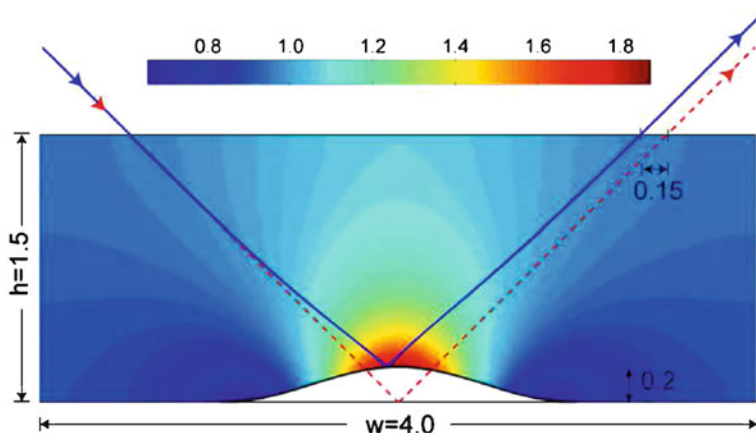


Fig. 5.38 Ray path onto a truly flat ground plain (*red dotted line*) and the actual ray path in the cloak (*solid blue line*). The *top color bar* indicates the value of the refractive index. The lateral shift between the two exit rays is comparable to the height of the cloaked cavity (Image courtesy of Baile Zhang/MIT)

Having dealt with three-dimensional invisibility cloaks, Zhang and co-workers next turned their attention to 2-D carpet devices. This latter work was first published online at the arXiv preprint data server on June 15, 2010, and describes how a carpet cloak must necessarily produce an anomalous lateral image shift that with careful observing should be detectable. The lateral shift comes about because the mirror substrate in any ground cloak must have some curvature in its profile (Fig. 5.38). Not only will the lateral shift reveal the presence of a carpet cloak, it also turns out that the shift is almost the same size as the height of the cavity being hidden⁶ – a good result, it would seem, for counter intelligence.

Transformation Optics

Although Michael Faraday pioneered the idea of electric and magnetic lines of force (recall Chap. 4, and Figs. 4.6 and 4.7), it was James Clerk Maxwell who brought together the mathematical equations that describe their behavior. Indeed, as we saw in Chap. 4, Maxwell's equations are the essential tools of classical optics – they literally describe how electromagnetic radiation behaves in the presence of matter. One of the remarkable characteristics of Maxwell's equations, however, is that they are invariant to changes in the coordinate system used to express them. What this means

⁶ A counter example to this situation is afforded by the calcite invisibility cloak described in Chap. 3. In this case, the birefringence properties of the calcite crystal are used to restore both the angle and position of the reflected ray (see lower diagram in Fig. 3.34). And, just to prove the point that it is difficult to poach from a poacher, the SMART calcite cloaking device, recall, was designed by Baile Zhang.

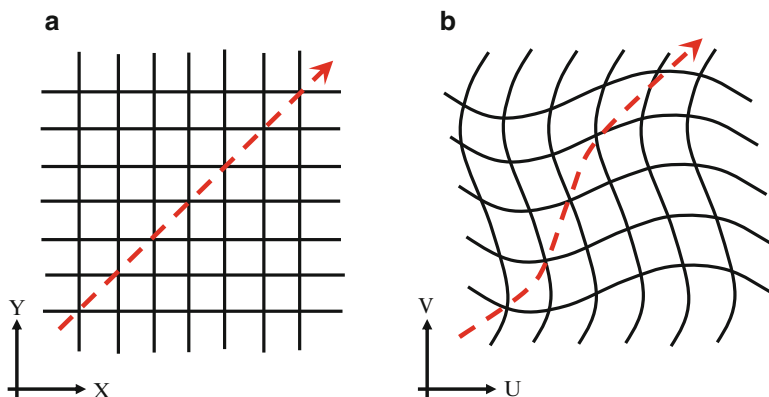


Fig. 5.39 The basis of transformation optics centers on being able to mathematically morph the distorted path and warped grid (B) on to the uniform grid (A). In mathematical speak, this means finding the transforms $u(x, y)$ and $v(x, y)$ (Image courtesy of J. Pendry and University College, London)

is that they look the same irrespective of the underlying geometry. Whether the equations are written in Cartesian coordinates (the standard x, y, z of three-dimensional space) or, say, in cylindrical coordinates⁷ they are identical, with the exception, that is, of the spatial variation of the ϵ and μ terms. The important point of all this is that it essentially tells us that the lines of force, which Maxwell's equations describe, are locked to the coordinate system, and if the coordinate system is distorted then the distortions will carry all the associated fields with it. It is this property of Maxwell's equations that allows transformation optics to work. The idea is that to make a particular light ray follow a certain distorted path, all we need to do is work out the coordinate distortions required to describe the path, apply these distortions to Maxwell's equations and then extract from the result the required expressions for ϵ and μ . This beautiful result was first made explicit and deterministic in a research paper published by Andrew Ward and John Pendry in 1996. It wasn't until the publication of the research papers by Ulf Leonhardt and John Pendry, David Schurig and David Smith in the June 23, 2006, issue of *Science*, however, that the field of transformation optics really came into its own (as indicated in Fig. 5.3).

Transformation optics provides a way, a computational way, of determining the spatial variation of ϵ and μ such that a light ray will follow a specifically chosen path. We can think of this in terms of rubber sheet geometry. Imagine a uniform grid of reference lines drawn upon a perfectly flat rubber membrane (Fig. 5.39a), with a specific light ray, moving along a straight line path, being projected onto it. Now, when the rubber sheet is stretched, warped and compacted both the light ray path and the grid line system are deformed in unison (Fig. 5.39a).

⁷ In this coordinate system the location of a point in space is described according to its radial distance r from a fixed central axis, the angle ϕ measured around the central axis and z the height of the zero point for the radius above a specific reference location.

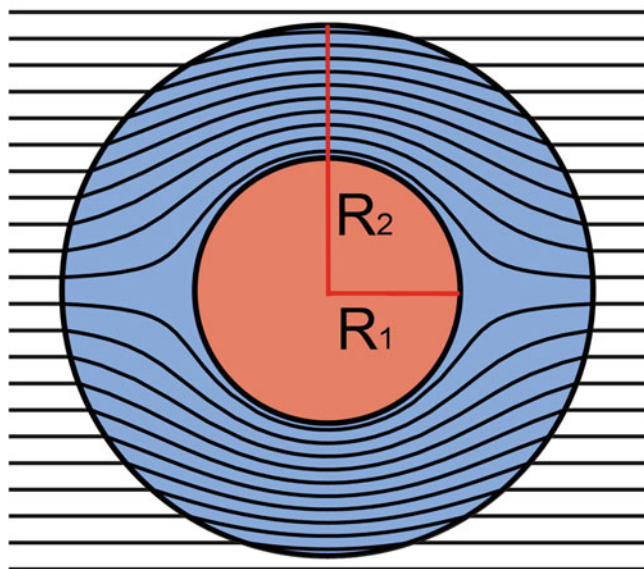


Fig. 5.40 Schematic image of ray paths required to produce an invisibility cloak. The invisibility region is that of interior to radius R_1 . The channeling of light begins at radius R_2 (Image courtesy of David Schurig)

Transformation optics works by finding the mathematical scheme that links the uniform grid pattern to the warped one. Once this transformation is known then Maxwell's equations can be solved to determine the spatial variation of ϵ and μ necessary to realize the warped light ray path. The final problem is then to take this spatial variation in the permittivity and permeability and fabricate it, point by point, within the interior of a metamaterial. In a very real sense transformation optics has a number of parallels with anamorphic designs (recall Chap. 1 and Fig. 1.9) where the distorted image (the rubber sheet) can be mapped back to a smooth mirror surface (the uniform grid).

In terms of designing an invisibility cloak, John Pendry, writing in the American Physical Society journal, *Physics*, for November 16, 2009, eloquently explains, "Design and construction of a cloak is in the nature of a grand challenge for optics. On the theoretical side, transformation optics provides a complete description for a cloak, but practical implementation requires mastery of metamaterial construction." therefore, based on the idea that in unfamiliar terrain it is initially better to walk rather than run, the first invisibility cloaks were designed to operate in just two dimension rather than three. Indeed, a two-dimensional invisibility cloak was considered, by way of an example of the transformation theory, by Pendry, Schurig and Smith in their 2006 *Science* paper. They write, "Our plan is to achieve concealment by cloaking the object with a metamaterial whose function is to deflect the rays that would have struck the object, guide them around the object, and return them to their original trajectory." In this manner the incident light rays are required to pass smoothly around a central circular cavity of radius R_1 (see Fig. 5.40).

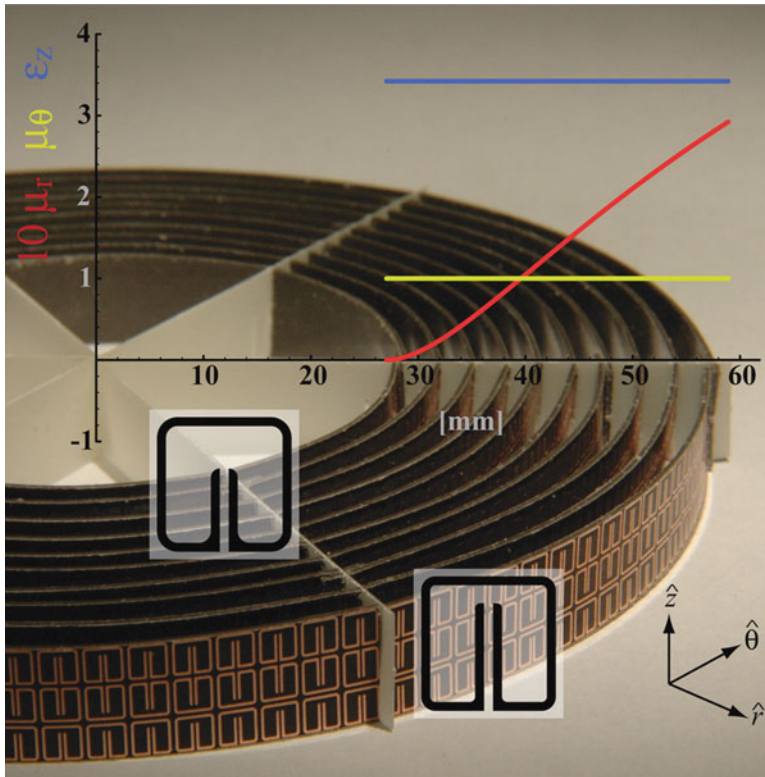


Fig. 5.41 This composite image indicates how $\epsilon(Z)$, $\mu(\phi)$ and $\mu(r)$ vary with radius r within the metamaterial disk. With these values the incident microwave radiation is channeled around the central circular cavity. Two split ring resonator profiles are shown as inset images (Image courtesy of Duke University)

Beyond a distance R_2 , the actual outer radius of the invisibility cloak, the light rays move along parallel paths, and the key point here is that to an outside observer there is no indication that an optical cavity has been traversed. In terms of rubber sheet geometry the cloak being produced can be imagined as the result of poking the tip of a cone through the uniform grid. As the tip of the cone is pushed through the sheet both a hole will open up and the nearby grid lines will expand and distort to smoothly accommodate it. These geometrical distortions can then, through transformation optics, be used to map out the required variation of the electrical permittivity (ϵ) and magnetic permeability (μ). The hole pushed into the rubber sheet by the expanding cone corresponds to the cloaked region that will be inaccessible to the incident light.

The variation of the electrical permittivity (ϵ) and magnetic permeability (μ) in the Duke University invisibility cloak, fabricated by Schurig and co-workers in late 2006, is shown in Fig. 5.41. Also illustrated in the figure are the profiles for the split-ring resonators used to create the required refractive index variation. Indeed, it is through changing the dimensions of the SRRs that the local permittivity and

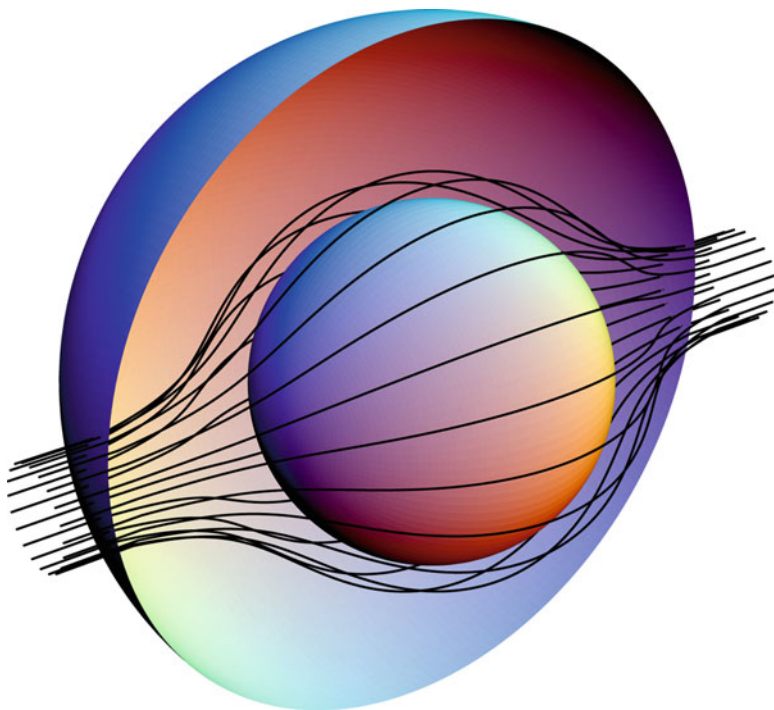


Fig. 5.42 Ray paths followed in a three-dimensional invisibility cloak. In this case the light rays are diverted around the small central sphere – the invisible cavity in which an object can be hidden from view. The region between the two spheres contains the all-important metamaterial substrate (Image courtesy of David Schurig)

permeability are controlled. The characteristic dimensions of the SRRs in each circular substrate are about 3 mm^2 .

Pendry, Schurig and Smith concluded their June 23, 2006, *Science* paper in triumphant tones, commenting that, “We have shown how electromagnetic fields can be dragged into almost any desired configuration. The distortion of the fields is represented as a coordinate transform, which is used to generate values of electrical permittivity and magnetic permeability ensuring that Maxwell’s equations are still satisfied. The New concept of metamaterials is invoked, making realization of these designs a practical possibility.” And, as we have seen, the realization of a two-dimensional cloaking device is exactly what David Schurig and co-workers at the Duke University and Imperial College labs announced online in their October 19, 2006, article in *Science Express*.

The ray paths required to realize a true three-dimensional invisibility cloak are illustrated in Fig. 5.42. In this case the permittivity and permeability will have to vary in a complex but now known and calculable way. To date no such metamaterial cloak has been fabricated, but it is surely just a matter of time, a decade maybe, before such a device is actually realized.

Harry Potter's Sleeves and Eyesight for the Blind

The ever-popular standby of many a science fiction adventure is the wormhole – a bizarre mathematical quirk of general relativity that technically allows a would-be space traveler to make short cuts between vast distances of space. In terms of Michio Kaku's list of impossibilities, the wormhole rates as a Type II to Type III borderline impossibility – that is, it is theoretically possible, but entirely unrealizable with present day (and indeed, near future) technology. So, while the generation of a spacetime wormhole in the laboratory is a project set in the far (infinitely?) distant future, the construction of an electromagnetic wormhole might take place somewhat sooner – although this being said, it won't likely happen any time soon.

An electromagnetic wormhole is not an actual general relativistic wormhole linking disparate regions of spacetime, but it is rather an extension of the invisibility cloaking idea. In an electromagnetic wormhole light is essentially transported invisibly from one location to another. This all sounds rather odd, so let us look a little deeper into the process.

The theory behind the formation of an electromagnetic wormhole was developed by mathematician Allan Greenleaf (University of Rochester) and co-workers in 2007. To the uninitiated, the mathematics is formidable, but the results of the analysis published in the October 12, 2007, issue of the journal *Physical Review Letters* are clear enough and they may also be very useful, having potential applications in medical research, the communications industry, and, of course, covert military surveillance.

The electromagnetic wormhole (EWH) is a twist on the invisibility cloak in that rather than being a spherical structure, a metamaterial cylinder or tube is fabricated instead (Fig. 5.43). While the outer shell of the tube is cloaked and invisible to an external observer, light can propagate down its central core. Indeed, a physical

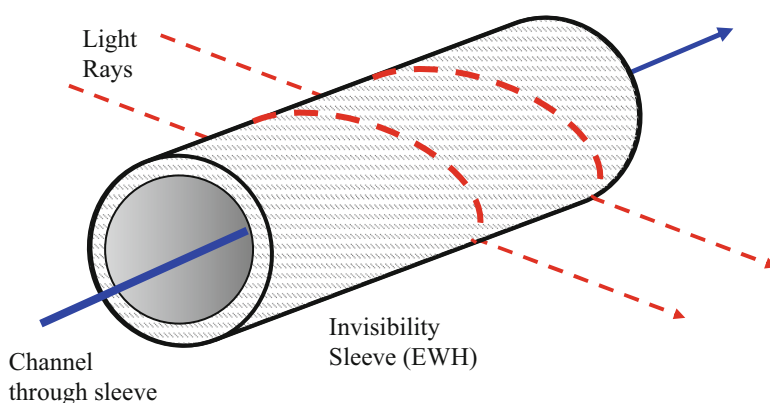


Fig. 5.43 Harry Potter's sleeve – or an electromagnetic worm hole (EWH). Here the cloaking is applied around the surface of a cylinder, leaving an invisible central core through which light and/or solid objects can pass

object can be moved invisibly down the cylinder (by any normal means – remember, no space warping actually takes place in an EWH). To an outside observer an object would disappear as it entered the EWH only to reappear again somewhere else at the other end of the cloaked cylinder. The vanishing and then reappearing at a different place trick of EWHs has prompted some researchers to describe them as a Harry Potter's sleeve.

One real-world application for EWHs that Greenleaf and co-authors have suggested concerns the application of medical surgery involving MRI imaging. The intense magnetic field generated by an MRI scanner normally precludes their use in surgical operations, since the surgeon's instruments distort the view. If the instruments are surrounded by an EWH metamaterial sleeve, however, then the MRI scanner would function without magnetic field distortion effects.

Greenleaf et al. have further suggested that EWHs might also be used to produce a true three-dimensional television. In this latter application a bundle of EWHs could be arranged (a bit like flowers in a vase) so that the light entering one end of the bundle, the image to be projected, reappears in a three-dimensional pattern at the other. While the bundles are invisible to the eye, the light exiting from the end of each EWH would appear as a free-floating pixel, arranged so as to produce a 3-D representation of the imaged object.

Where EWH devices might find their ultimate application is in providing eyesight to the otherwise totally blind. Indeed, without something like a EWH viewing port, an observer situated within an invisibility cloak is totally isolated from the outside world – no light gets in and no light gets out of the cloaked region. If appropriately placed and sized EWH viewing tunnels are smoothly melded onto an invisibility cloak, however, then a view of the outside world could be realized. The process works both ways, of course, so the design would have to be carefully crafted in order to prevent an external viewer from seeing into the interior of the cloaked cavity. Here, then, with an EWH sleeve do we have the potential for constructing the ultimate invisibility device; not only hidden from external view, the cloaked operative would no longer sit in visual isolation but could in principle move freely around and view the outside world from their hidden optical cavity – the manifestation of magic indeed.

Where Next for Metamaterials?

Nanotechnology and metamaterials are going to change the world we live in. While medicinal and pharmaceutical advancements have already been made, nanotechnology will ultimately deliver major advancements in electronic miniaturization, data storage, increased computing speed, faster communications, rapid light manipulation devices and, of course, invisibility cloaks. This being said, science correspondent Erin Biba writing in the August 2010 issue of *Wired Magazine* suggests that we are perhaps 20–200 years away from the practical realization of optical invisibility cloaks. She saliently points out, however, that, “at least we can see it coming.”

The potential applications for metamaterials are manifold, but perhaps the most exciting discoveries will be made in the field of pure physics research. Ivory-towered though such work might be, it will be at the forefront of scientific creativity, opening the metaphorical doors to laboratory-based research on analog devices describing, for example, astrophysical black holes (see Appendix 1), the early universe and light propagation in highly curved space.

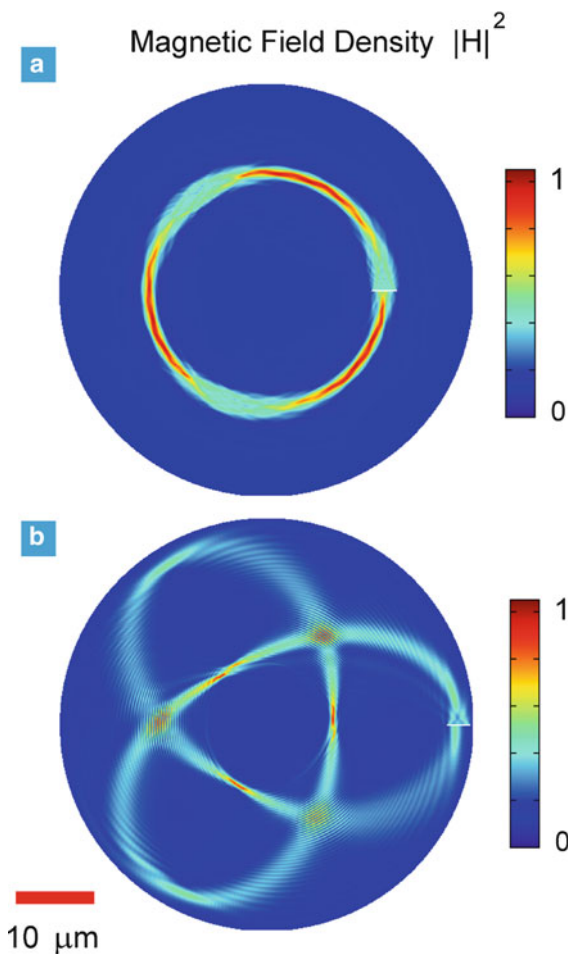
Optical Dynamics and Experimental Cosmology

The Berkeley research group, headed up by Xiang Zhang, has recently begun investigating the possibility of mimicking celestial dynamics with metamaterial analogs. Writing in an article published online by the journal *Nature* on July 20, 2009, the Berkeley group under lead author Dentcho Genov (Louisiana Tech University) showed that a whole series of classical, special and general relativistic dynamical problems might be investigated with metamaterial devices. Indeed, the group draws upon a parallel that exists between the least-action principle of mechanics and Fermat's least-travel time principle in optics (recall from Chap. 2) to produce a series of detailed simulation calculations. Genov and co-authors also note that one area where specially fabricated metamaterials, configurations they describe as continuous index photon traps (CIPTs), might come into their own is in the study of orbital stability and chaos (Fig. 5.44). Indeed, chaotic systems abound in nature, encompassing such diverse topics as predatory population dynamics and the time variation in the orbit characteristics of solar system objects. Typically extensive, time consuming numerical calculations are required to investigate the possible behaviors of a chaotic system, but CIPTs offer the potential for the rapid analog evaluation of such systems in a laboratory setting.

The field of experimental optical dynamics is barely beyond the threshold of dreams at the present time, but the future for such experimentation looks remarkably bright. Igor Smolyaninov (University of Maryland) has perhaps looked the deepest into the possible uses of metamaterials as analogs for extreme physical environments. Posting his research thoughts in a series of papers uploaded to the arXiv server during the summer of 2010, Smolyaninov, along with collaborators Alexei Smolyaninov (University of Maryland) and Evgenii Narimanov (Purdue University), have outlined ways in which the "metamaterial landscape" might be manipulated in order to mimic those conditions thought to apply at the boundaries between adjacent universes in the multiverse (recall Chap. 1).

The physics of five-dimensional spacetime might also be amenable to study through appropriately constructed metamaterial analogs, as well as the 'foamy' spacetime that some theories of quantum gravity require. Not only this, but different spatial topologies might also be amenable to investigation through the use of special split ring resonator designs. Indeed, Chih-Wei Chang and co-workers at the University of California at Berkeley describe in the 3 December 2010 issue of *Physical Review Letters* an optical metamaterial with a twist. In this case it is a

Fig. 5.44 Simulation of infrared light propagation in an air to GaInAsP (gallium, indium, arsenide, phosphide) high refractive index photon trap. (a) Circular orbit with a period of 2π . (b) Rosetta orbit of period 4π . The *sidebars* show the variation in the magnetic field density relative to the incident beam (having $|H|^2 = 1$), while the *scale bar* in the lower left-hand corner gives a 10 μm measure (Image courtesy of Dr. D. Genov, Louisiana Tech University, Ruston, Louisiana)



topological twist that they induce (rather than a physical substrate twist) through the use of what they call metamolecular trimers (a modified form of split ring resonator). One of the trimer configurations considered by Chang and colleagues (Fig. 5.45) can (in principle) twist the geometry, as seen by a light ray, to mimic that of the one-sided Möbius strip.⁸

Incredibly, it may even be possible to perform laboratory-based experiments on the physical processes that took place during and immediately following the Big

⁸ A Möbius strip can be made from a rectangular piece of paper. The trick is to join the two shorter edges of the paper band together but adding a half-twist into it before making the ends fast. If you run your finger around the Möbius strip, so produced, you will find that it has only one edge and only one surface.

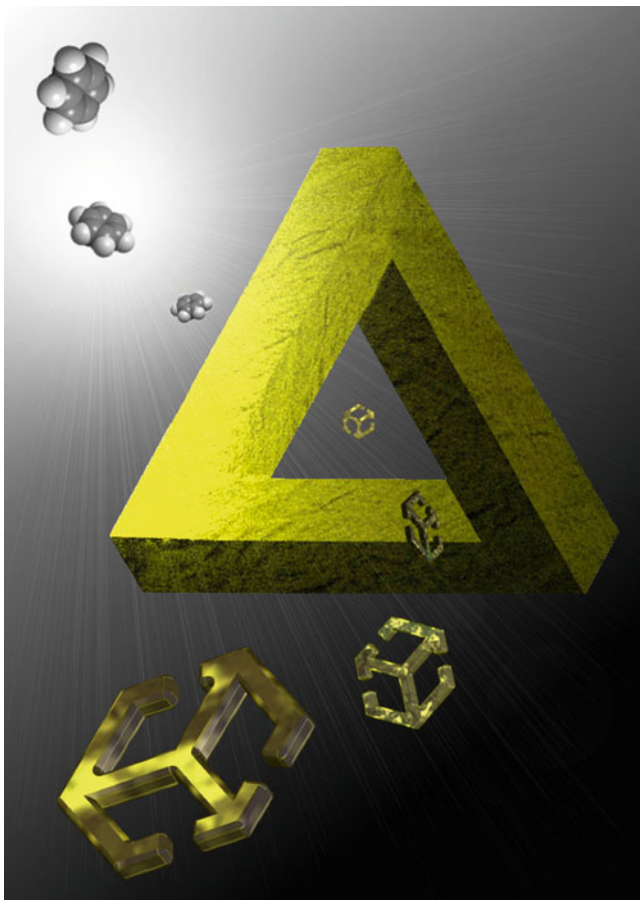


Fig. 5.45 The Penrose tri-bar (recall Fig. 1.7) has a Möbius symmetry—meaning that it is a one-sided surface. The metamolecular trimers that can mimic this geometry are shown to the lower center and lower left of the image (Image courtesy of University of California, Berkeley Labs)

Bang. This is truly an astounding possibility, making the inherently invisible formative moments of the universe amenable to direct study. Indeed, uploading a research paper to the arXiv web server on 4 April, 2011 Igor Smolyaninov and Yu-Ju Hung (University of Maryland) described a metamaterial “toy model” designed to emulate the flow of time in the earliest stages of the Big Bang. Remarkably, experiments conducted with the model have enabled Smolyaninov and Hung to look at such fundamental issues as the second law of thermodynamics, which states that disorder must increase over time, and the cosmological arrow of time, which provides a direction of evolution for universal expansion.

Not only has Smolyaninov considered how experiments with metamaterials might provide direct analog models of the early universe, he has also looked at ways

in which the curvature of spacetime itself might be studied. Indeed, posting his ideas on September 28, 2010, Smolyaninov looked at the possibility of fabricating metamaterials to investigate the warp drive model developed by Miguel Alcubierre (recall Chap. 4 and Fig. 4.19). The so-called metric that defines the warp drive spacetime has been extensively studied in recent years, and many seemingly fatal problems to its application have been raised. Possibly the most serious problem associated with the Alcubierre scheme is that it violates the conservation of energy law – one of the most fundamental laws in all of physics. It is also unstable when quantum mechanical effects are introduced. As always, there are exotic ways of getting around such problems, but they invoke the existence and use of such hypothetical entities as negative energy matter, and even tachyonic matter (the latter being composed of particles that always move at speeds greater than that of light⁹). All the above being said, Smolyaninov argues that it might just be possible to construct a metamaterial analog that emulates Alcubierre's metric, and that this might allow for a gradually accelerating warp drive model to be studied directly.

Working with more traditional materials, Vincent Schultheiss (Max Plank Institute for the Science of Light at Erlangen in Germany) and co-workers have recently performed experiments to investigate the behavior of light constrained to move upon highly curved surfaces. Writing in the journal *Physical Review Letters* for October 1, 2010, Schultheiss and colleagues describe how they have produced ultra-smooth glass and aluminum surfaces corresponding to a sphere and a negatively curved hyperbolic surface (Fig. 5.46); these shapes act as three-dimensional analogs to the warped space associated with strong gravitational fields (the sphere) and some hybrid cosmological models (the hyperbolic surface). The optical surfaces were fabricated with a specialized diamond polishing machine, and in the case of the aluminum surface coated with a thin layer of oil. Laser light was then fed tangentially into the surfaces of the two substrates causing it to become trapped. The manner in which the light propagated along and around the substrate surfaces was then analyzed. The experiments, still in the early stages of development, will eventually, it is hoped, help visualize how light travels in regions of high spatial curvature, and also allow for the construction of analog models relating to general relativity.

While the fabrication of metamaterials having the characteristics necessary to study the strange physics of warp drives, the early universe and even the multiverse are far from being realized, the exciting point is that we are now essentially quibbling about manufacturing processes rather than abstruse theoretical details. In the future the present-day metamaterial revolution might well precipitate the birth of whole new sciences – the science of experimental cosmology, for example. These are perhaps naively optimistic dreams, but dreams are often the precursors of reality.

⁹ Tachyons are entirely possible, but we have at present no idea if they actually exist. While they move faster than the speed of light they do not contravene Einstein's special relativity, which only postulates that the speed of light is a limiting speed. No sub-light speed particle can be accelerated to the speed of light, and likewise, no tachyon can be slowed down to subluminal speeds.

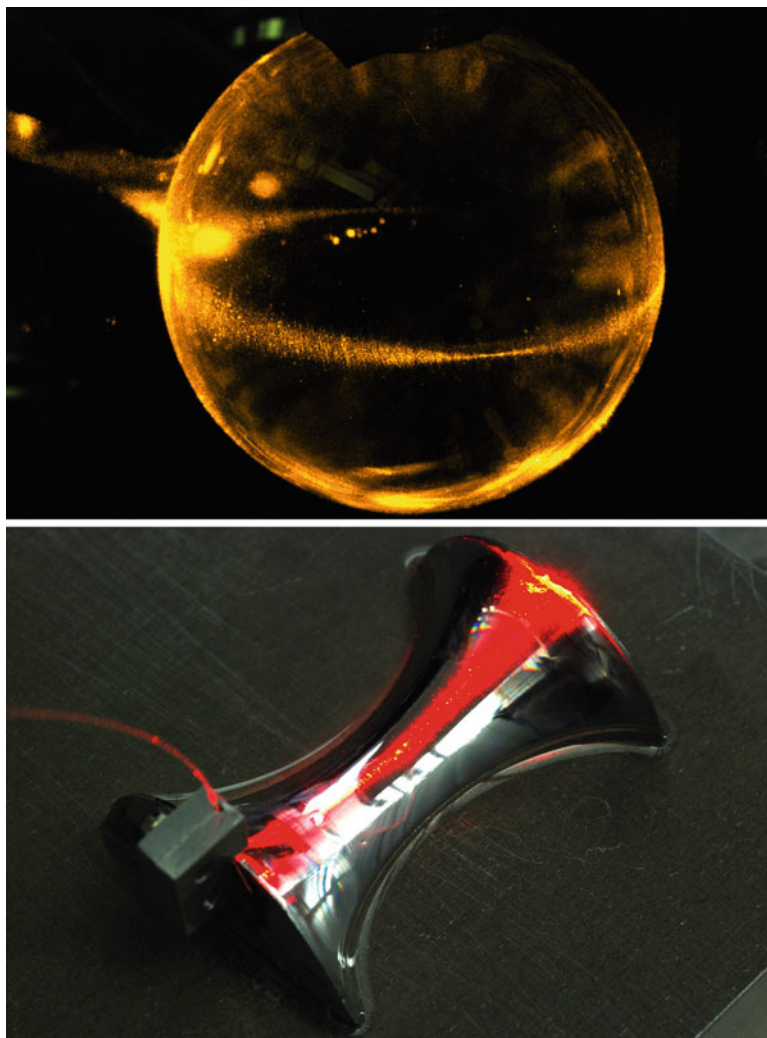


Fig. 5.46 (Top) Light trapped upon the ultra-smooth surface of a 2.5 cm diameter glass sphere, as studied at the Max Planck Institute for the Science of Light. Interestingly, it can be seen that the incident laser beam is refocused after traveling half way around the circumference of the sphere – this characteristic is similar to that expected for Maxwell’s fish-eye lens (recall Fig. 5.17). (Bottom) Light propagation over the surface of a negatively curved substrate. In this case the light beams spread outward around the surface (Images provided by Sascha Batz, University of Erlangen-Nuremberg)

Out of Space and Time

“Would it not be amazing if, in an age of ubiquitous surveillance, a method could be devised to somehow remove a piece of history that would otherwise be recorded by a surveying camera?” These words are not taken from a science fiction story but

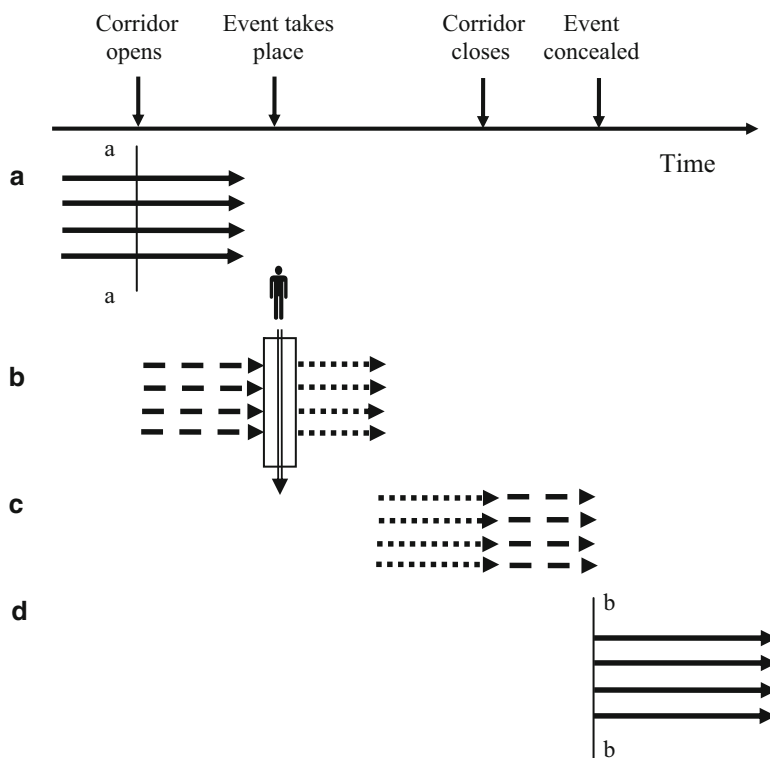


Fig. 5.47 The operational principle of a spacetime cloak (STC). (a) A light ray enters the STC and at time a-a the various light rays are split into two components, with the leading part having a higher speed than the trailing part. (b) A time gap corridor is opened up between the leading and trailing rays allowing an object to be moved across the light stream. (c) The time gap closes and the speed of the leading rays is reduced while those of the trailing rays increased. (d) From time b-b onwards the light rays are recombined and all evidence of the time gap corridor has been removed, with the observer seeing a continuous and uninterrupted flow of light

rather constitute the opening sentence of an academic paper published in the *Journal of Optics* on February 2, 2011. The paper was written by Martin McCall (Imperial College London) and co-workers and concerns the possible construction of a spacetime cloak (STC) that hides events rather than objects.

In its operation an STC creates a temporary corridor through which information, or even matter, could be transported, and then hides the event from external view. The trick to making this work is to use a metamaterial substrate to alternatively speed up and then slow down light rays. Figure 5.47 illustrates the idea. Upon entering the STC a light ray is effectively chopped into two components. The front, leading part of the light ray is sped up, while the second trailing part is slowed down. This chopping process creates a gap, or corridor, in the light ray bundle through which something could be quickly moved, or transmitted. Once the message or object has traversed the gap, the STC then acts to slow down the leading part of the

light ray and speed up the trailing part – allowing them to exit the STC substrate with absolutely no trace of the temporary corridor having been created. Indeed, an external observer will record a continuous, uninterrupted flow of light.

In a lecture delivered to the Physics Department at Imperial College London on March 2, 2011, McCall presented a fictional, but illustrative, scenario in which an STC could allow a thief (dare we say a light-fingered one?) to enter a continuously monitored room containing a safe and thereafter proceed to rob its content without leaving any incriminating visual evidence.¹⁰ The crime would proceed by setting up a time gap corridor with an STC. Then, while the time corridor was open, the thief could rush across the room, crack open the safe, take out its content, close the safe door again, and finally exit the room. Provided that the thief was quick acting and left the room before the STC corridor closed the security personnel monitoring the room would be none the wiser that the contents of the safe had been stolen. To their reckoning no one entered the room, the safe door was never opened (because the clever thief closed it after removing the contents) and its treasure was safe inside. A perfect crime.

Although presented in terms of an editing-out of events device, there are, in fact, many real world problems that an STC could solve – specifically in the areas of signal processing and high-speed computing. When, however, is the first STC likely to appear? Well, this is likely to be some time in the future – perhaps within the next several decades. Interestingly, however, McCall and co-workers comment in their publication that a proof-of-concept experiment could be conducted with present-day technology. Indeed, the basic idea could be tested with nonlinear fiber optic cables in much the same way that analog black holes and Hawking radiation (see Appendix 1) are already being studied.

McCall et al. conclude their research article (echoing one of the major recurring themes of this book) by noting that, “whilst the camera is said to never lie, it does not always record the whole truth.” Indeed, with the STC, once developed and fabricated, we reach the ultimate in the control of light images. Not only might an object, or parts of an object, be hidden from our direct gaze, or morphed to appear like something else, elements of its dynamical history might also be edited out from direct view as well. Magritte’s “Treachery of Images” (Fig. 1) will then have been taken to its ultimate space and time extreme.

¹⁰ A streamed video of the 60-min lecture by McCall entitled *Lies, damned lies and tricking the light fantastic*, can be downloaded from the Imperial College London media library.

Afterword

Science and science fiction writer Arthur C. Clarke spent many years of his fruitful life thinking and writing about the future. He sadly passed away in 2008, just as the initial phase of the metamaterial and invisibility cloaking revolution was in its ascendancy. We can be sure, however, that he followed the early beginnings of this new research field with interest. And, while we do not know Clarke's actual thoughts on the technological innovations it seems reasonable to suppose that he would have marveled at the new developments. Indeed, the physics of invisibility, rightly now so-called, in less than the span of one decade, has lived up to those prophetic words Clarke wrote in *Profiles of the Future* (published in 1962), that, "Any sufficiently advanced technology is indistinguishable from magic."

The magic $\Leftrightarrow$ advanced technology synonym is often referred to as Clarke's third law. His first law is, "When a distinguished but elderly scientist states that something is possible, he is almost certainly right. When he states that something is impossible, he is probably wrong." The second law explains that, "The only way to discover the limits of the possible is go beyond them into the impossible." All three of these laws have applied, at one time or another, to the story of metamaterials and the development of invisibility cloaks, as we have seen in the preceding chapters. The physical magic of invisibility cloaking, however, has now been broadcast and its thaumaturgical pulse is thrumming with innovation, literally tripping the light fantastic, and rushing headlong into our future. For all this, the physics of invisibility is still a young science; what it might produce when fully matured will no doubt truly amaze us. Indeed, we will probably not believe our eyes, and it seems only reasonable to speculate that the possible applications of this new technology are limited only by our imaginations – the magic of metamaterials and invisibility is manifest.

*Something hidden. Go and find it.
Go back and look behind the Ranges -
Something lost behind the Ranges.
Lost and waiting for you. Go.*

Rudyard Kipling – The Explorer

Appendix 1

Black Holes in the Laboratory

It seems too far fetched to even consider – study a black hole in the laboratory? This, given the potential danger of such matter destroying objects, is surely no laboratory any sane scientist would want to be experimenting in? But fear not, these black holes are entirely harmless, impotent doppelgangers of their astrophysical cousins. Remarkably, we actually find ourselves living in the age when the first analog black holes are being made in laboratories around the globe. Indeed, researchers are now actively investigating the physics of these tamed gravitational monsters, and specifically they have begun to study the emission characteristics of Hawking radiation.

First proposed by Stephen Hawking in the 1970s the diaphanous radiation named in his honor is produced at the event horizon surrounding the central singularity of a black hole. Building upon the quantum mechanical idea that space, even a vacuum, is full of virtual particle-anti-particle pairs rapidly popping into and then out of existence, Hawking argued that at the event horizon one of the virtual particles might end up being captured by the black hole with its companion escaping into space. This fitful capture and escape process on the very threshold of a black hole's one-way, trap-door boundary should result in the outward propagating of radiation – Hawking radiation. One result of all this complex and heady physics is that black holes will slowly, ever-so-slowly, evaporate away – even the gargantuan black holes at the centers of galaxies will eventually fizzle to nothingness. Astronomers have not as yet detected Hawking radiation in space, but this is where the new analogs, located conveniently within terrestrial laboratories, might pave the way towards its eventual detection. Indeed, some would say it has already been demonstrated.

It turns out that there are a number of ways of constructing perfectly safe black hole analogs. The trick is to parameterize gravity by some other physical variable, and then reproduce the conditions expected to apply in the region of an event horizon. University of British Columbia, Canada, physicist William Unruh, has likened the processes that occur close to an event horizon as being similar to those experienced by a fish trying to swim upstream, away from a waterfall or weir – where the boundary over which the water falls represents the event horizon. The point that Unruh is making is that light (the fish) propagating in a moving medium (the stream)

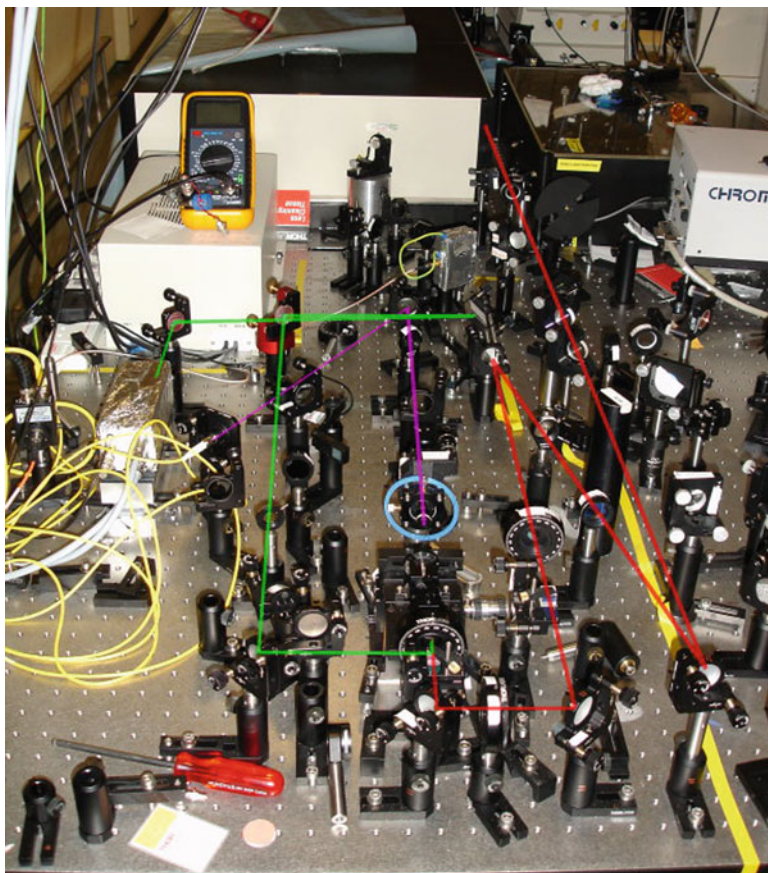


Fig. A1.1 Optical table arrangement for the fiber optical analog of a black hole. The all important optical fiber is wound into a circular loop, visible in the center of the picture (Image courtesy of the University of St. Andrews)

behaves in the same way as light propagating in the curved spacetime close to the region of an event horizon.

Using Unruh's swimming fish analog Ulf Leonhardt and co-workers at the University of St. Andrews have developed a special optical fiber analog for the event horizon surrounding a black hole (Fig. A1.1). The St. Andrews group made use of a so-called microstructured optical fiber which has a hexagonally shaped core composed of numerous air-filled cylinders. Writing in the March 7, 2008, issue of the journal *Science*, Thomas Philbin, along with Leonhardt and co-workers, explain that the special characteristic of such optical fibers is that their refractive index changes slightly as a pulse of light propagates along it – the so-called Kerr media. The crucial point here is that the refractive index perturbation within the fiber can be thought of as an analog of the gravitational field in the region of an event horizon.

The St. Andrews black hole analog is activated by first sending an energetic pulse of laser light down the fiber. This initial, soliton pulse has a constant amplitude and

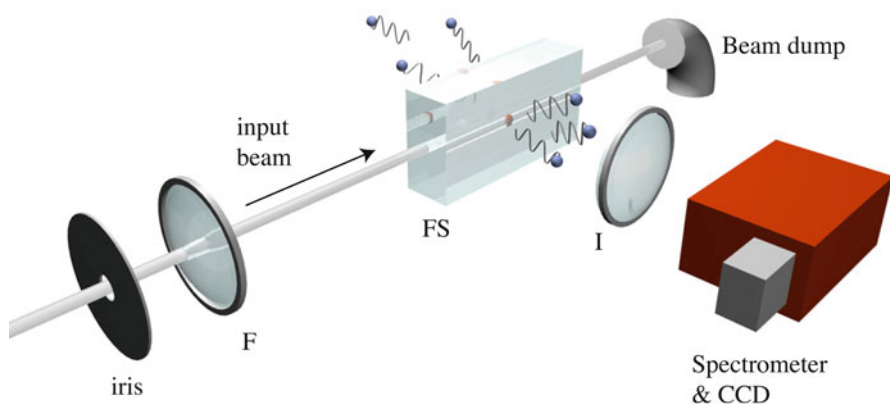


Fig. A1.2 Schematic layout of the experimental arrangement used to successfully detect analog Hawking radiation. FS is the fused silica glass block target, and the analog Hawking radiation is captured through lens I with the spectrometer and CCD recording instruments (Image courtesy of Daniele Faccio, Heriot-Watt University, Edinburgh, Scotland)

wave profile and as it races along the fiber it very slightly increases the refractive index behind it – importantly, this variation in the refractive index moves along the fiber with the laser pulse. We can think of the soliton pulse as being a moving event horizon (or waterfall boundary in Unruh’s picture). The trick now is to send a second laser pulse down the fiber to chase after the initial soliton pulse (this will be the swimming fish). The second pulse has a different wavelength to the soliton and therefore, and because of the properties of the optical fiber will propagate at a faster speed, moving forward so as to catch-up with the soliton pulse. Because of the increased refractive index behind the soliton, however, the second chasing pulse is slowed sufficiently that it never actually overtakes the leading soliton,¹ and this is the event horizon analog. A black hole has been safely created out of light, within a strand of glass, in the laboratory – incredible!

The St. Andrews researchers have verified that their optical fiber black hole analog works, but have not as yet been able to detect any Hawking radiation from the region behind the soliton pulse. The problem is a highly technical one since to detect any such radiation the light of the soliton pulse itself must first be subtracted out. Research continues.

While the St. Andrews work proceeds, the first successful detection of Hawking radiation in an analog experiment was announced online, at the arXiv data base, for September 23, 2010. Using an ultrashort laser pulse technique, Francesco Belgiorno (Universita degli Studi di Milano, Italy) and co-workers were able to generate and record Hawking radiation in a fused silica glass block (Fig. A1.2). Working with a laser system at Insubria University in Como, Italy, each light pulse in their experiment

¹ Technically under the conditions just described a white hole has been produced. White holes are technically the time reversed counterparts of black holes and they act as spontaneous matter generating sources. The point, however, is that both objects are surrounded by an event horizon.

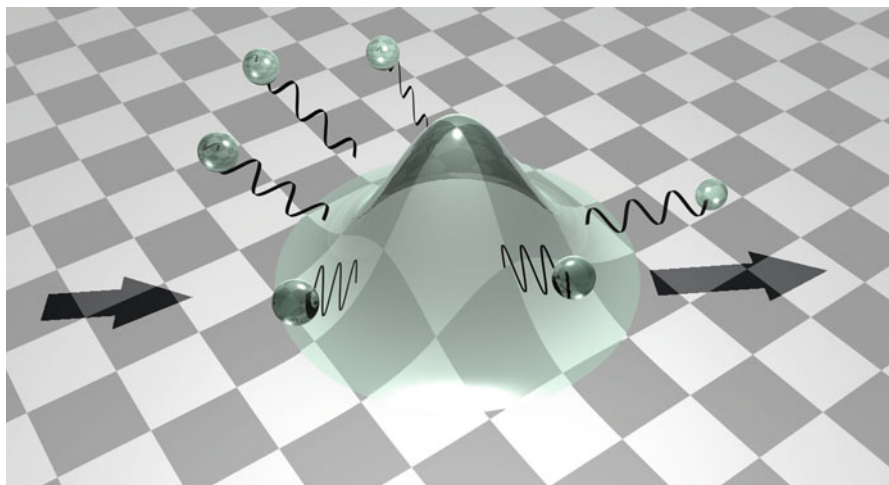


Fig. A1.3 Hawking radiation in association with a soliton RIP. The constant profile, soliton, pulse moves in the direction indicated by the *solid arrows*, creating analog black hole (*leading*) and white hole (*trailing*) event horizons as it moves through the optical medium. Photons produced at the soliton boundary will be analogs of the Hawking radiation produced at the event horizon surrounding a black hole (Image by and courtesy of Daniele Faccio, Heriot-Watt University, Edinburgh, Scotland)

was just 1-ps (one million millionth of a second) in duration, and the beam was fired directly at a 2-cm-long fused silica (FS in Fig. A1.2) glass target.

As each intense pulse of light propagates through the glass block its electric field briefly alters the refractive index, and a so-called refractive index perturbation (RIP) is generated. Two event horizons can be associated with each RIP, with the leading edge being analogous to a black hole event horizon and the trailing edge analogous to a white hole event horizon (Fig. A1.3). The monitoring equipment is carefully positioned so as to detect only radiation that is emitted from the fused silica target at right angles to the direction in which the RIPs are propagating. In this way any contamination of the Hawking radiation from other light-emitting phenomena is minimized. Belgiorno et al. write in their discovery article that, “We report experimental evidence of photon emission that on the one hand bears the characteristics of Hawking radiation and on the other is distinguishable and thus separate from the known photon emission mechanism.” To this they add in modest and almost understated tone, “We therefore interpret the observed emission as an indication of Hawking radiation induced by the analogue event horizon.”

What an incredible experimental result this truly is, and 36 years after Hawking first suggested that such radiation should be emitted in the region of a black hole event horizon his perceptive ideas have now been experimentally vindicated in analog form.

Appendix 2

Atmospheric Refraction, Extinction and Scattering

Small density variations in the many layers that make up Earth's atmosphere result in light rays being refracted and scattered as they propagate towards ground level. These effects are at their greatest when an object (star, planet, the Sun, or Moon) is close to the horizon. Indeed, all bodies appear higher in the sky due to atmospheric refraction than they are geometrically, and all objects close to the horizon will appear fainter than they really are due to the greater geometrical depth of the atmosphere that they are seen through. Indeed, the geometric depth of the atmosphere in the direction of the horizon is ten times greater than the zenithal geometric depth.

Take a look at Fig. 3.18 one more time. The atmosphere is approximated in the diagram as a series of m horizontal slabs, each of which has a slightly different refractive index running from n_0 to n_m . Given this approximation of Earth's atmosphere a remarkable result comes about, in that if Snell's law is applied at the boundaries between each of the layers it turns out that the total atmospheric refraction effect is determined by the refractive index at the very first and the very last layer.

Accordingly, if Z_0 is the measured zenith angle of a star at Earth's surface and Z_m is the true (geometrical and no atmosphere) zenith angle, then the following Snell's law sequence can be constructed:

$$n_0 \sin(Z_0) = n_1 \sin(Z_1) = n_2 \sin(Z_2) = \dots = n_{m-1} \sin(Z_{m-1}) = n_m \sin(Z_m)$$

From this we can see that no matter how many layers are assumed to be in place, that is, m can be very large, we always end up with the same result that: $n_0 \sin(Z_0) = n_m \sin(Z_m)$. At the top of the atmosphere the density is so low that the approximation $n_m = 1.0$, corresponding to a vacuum, will hold to a very good approximation. At ground level we can take a typical value of $n_0 = 1.000293$ (that of air). We now have enough information to determine the atmospheric correction term $\delta = Z_m - Z_0$, for any specific zenith angle Z_m . The results are shown in Fig. A2.1.

Using Fig. A2.1 as our guide it can be seen that when the true (geometrical) zenith angle $Z_m \approx 90^\circ$ then the atmospheric correction term is $\delta \sim 0.5^\circ$, and accordingly, the apparent (observed) zenith angle will be $Z_0 \approx 89.5^\circ$ – in other words, objects appear about 0.5° higher in the sky than they would if there were no atmosphere.

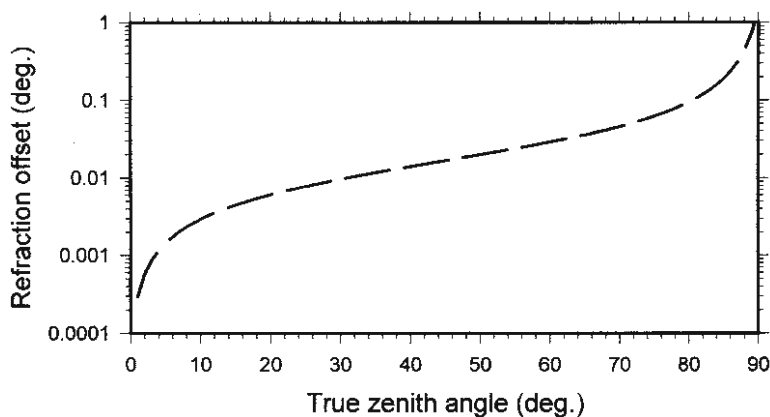


Fig. A2.1 The atmospheric correction term δ versus the true zenith angle Z_m

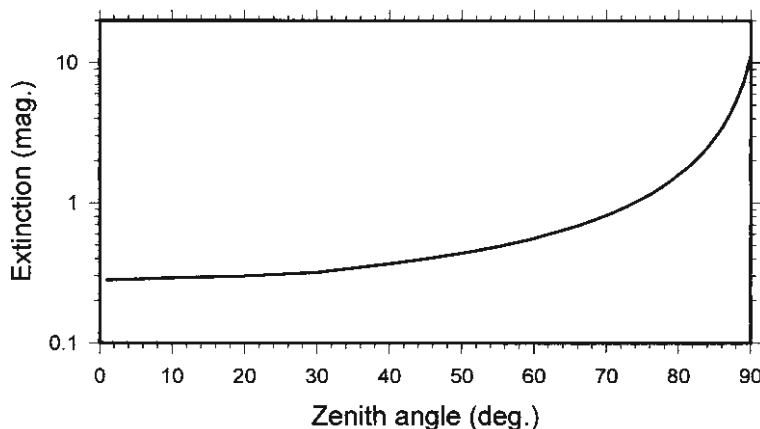


Fig. A2.2 Atmospheric extinction versus the true zenith angle Z_m

As the zenith angle decreases towards zero (i.e., towards overhead), however, so the refraction correction term δ becomes smaller and smaller and Z_m all but equals Z_0 .

Correcting for atmospheric extinction, or dimming, is a non-trivial problem, but Daniel Green has published a useful series of correction tables in the *International Comet Quarterly* (www.icq.eps.harvard.edu/ICQExtinct.html). Figure A2.2 shows the correction in magnitudes against zenith angle for an observer located at sea level. The correction term is to be added to the brightness (that is, apparent magnitude) of the object being observed. As with the refraction correction term, the dimming is smallest towards overhead but increases dramatically close to the horizon.

Indeed, on the horizon the extinction due to the atmosphere amounts to some 11 magnitudes. In this manner, for example, given that the Sun, in the zenith, has an

apparent magnitude of -27 , so on the horizon it will appear as magnitude -16 , which is just a little bit brighter than the full Moon. This is why it is safer to look at the Sun with the unaided eye when it is on the horizon but not when it is positioned high in the sky. Likewise, the full Moon (with an apparent magnitude of -13 in the zenith) will appear at magnitude -2 when it is located close to the horizon.

Not only are objects higher in the sky (by about $\frac{1}{2}$ of a degree) and fainter when viewed close to the horizon, they also look redder. This, as we described in Chap. 3 with respect to the Green Flash, is due to the blue wavelengths of light being preferentially scattered out of the direct line of sight through interactions with atmospheric molecules and dust.

One further odd horizon phenomenon should be mentioned here, but it is a phenomenon due to how the brain interprets images rather than an effect caused by Earth's atmosphere. The so-called Moon illusion was noticed many thousands of years ago and relates to the fact that the Moon appears to be much larger in angular size when it is seen close to the horizon than when it is seen high in the sky. This effect is not due to refraction, which merely displaces the whole lunar image upwards, nor is it related to the brightness and apparent color change of the Moon with zenith angle. Indeed, the Moon illusion is an entirely brain-produced phenomenon, but psychologists are still undecided as to why this effect comes about. When all is said and done, however, it is an effect that can be readily seen for oneself – the size of the Moon is indeed, in the eye of the beholder.

Appendix 3

How Big Is a Photon?

The following experiment is taken from David Deutsch's lectures on quantum computing. It is a nice thought experiment that gets to the core of the question, but the reader might like to think about some of the practical problems that would be encountered in trying to implement the experiment in practice.

Imagine taking two red laser pointers – just like the ones used to highlight text and pictures on a display screen during a talk or class. The fact that they are red laser pointers tells us that the light they emit has a wavelength of 6.5×10^{-7} m (see Fig. 4.12). The laser pointer used by the author has a power rating of 5 mW and this means that when in use the pointer emits an impressive 10^{16} (that is, 10,000 million million) photons per second – call this number R , corresponding to the rate of photon emission.

Let us further imagine performing an experiment with the two laser pointers: set them up so that the beams cross one another at right angles, with one beam being horizontal and the other vertical (Fig. A3.1). The two beams pass straight through each other at their intersection point and hit their respective target screens as two small red circles. When the author watched two such beams in a darkened room for one minute there was absolutely no sign, to the unaided eye at least, of any photons being scattered outside of their red circles on the screens. With this experimental result, no photon scattering observed in 1 min (and remember each laser emits 10,000 million million photons per second), we can estimate the approximate size of the photon.

Let each photon be a sphere of diameter d . Now, the number of photons per unit meter length of the laser beam will be given by the photon rate R divided by the speed of light c . Accordingly, the number of photons per unit area of the laser beam (the so-called area density A) is $A = R/cW$, where W is the width of the beam (approximately 1 mm). With this result in place, the probability of a photon crossing the beam striking another photon and then being scattered will be the cross-section area of the photon times the area density: $(\pi d^2/4) A$. Now, in time T the number of photons crossing one of the laser beams will be the product RT . The fact that no scattering is observed in time T means that we have a final result $RT(\pi d^2/4)$

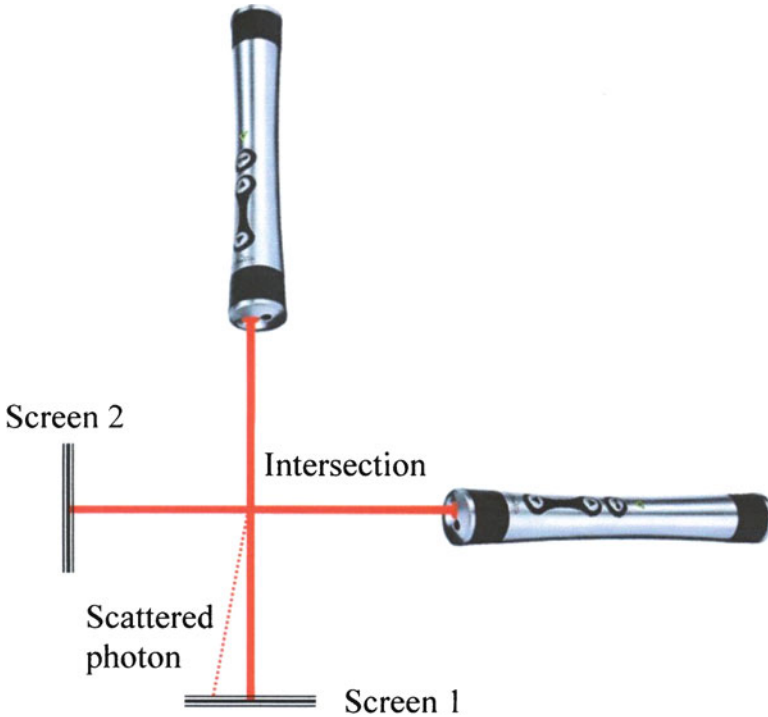


Fig. A3.1 Schematic arrangement of the laser pointers in an experiment for estimating the size of a photon

$A < 1$ – that is, the probability of a photon being scattering is less than 1 (i.e., no such event was observed) when the time interval T is less than 1 min (the specific observation time).

We can now use the known values for R , T and A to estimate the size of the photon. Accordingly, we find $d < 10^{-14}$ m, which is a very small size. Indeed, it is difficult to appreciate just how small the photon is in everyday terms, but if one imagines expanding the little red circle of the laser pointer's light to the same diameter as Earth's orbit around the Sun (a diameter of about 300 million kilometers), then according to our estimate the individual photons would be no larger than about 1 mm in diameter.

Index

A

Abbott, E., 11
Aeneid, 19
Aether, 42, 43, 120, 121, 123–127
Alcubierre, M., 128, 184
Alhazen, 33–34
Ampere, A., 112
Ampere's law, 112, 113, 122, 141
Anamorphic, 25, 26, 82, 84, 176
Anderson, H.C., 15
And He Built a Crooked House, 9
Aristotle, 32
ART. *See* Augmented reality technology (ART)
Asteroid (216 Kleopatra), 31
Astounding Stories of Super-Science, 3
Astronomy Nova, 37
Augmented reality technology (ART), 67, 71

B

Bacon, R., 34
Barentz, W., 86
Bennett, A., 4
Big Bang, 29, 183
Blackbody radiator, 160
Black hole, 11–14, 160, 161, 181, 187
Blaikie, R., 151
Blind spot, 27, 28
The Blind Watchmaker, 27
Blue screen, 21
Boyle, R., 53, 55
Bradley, J., 45–46
Brahe, T., 37

C

Calcite
 invisibility cloak, 97–99, 174
 prism, 97–99
Camouflage, 4, 8, 9, 19, 21–23, 28, 63, 70, 89–91, 133
Cassini, G.D., 43, 44
Cats' eyes, 64–67, 91–93
The Cave of Horror, 3
CERN, 11
Cherry-Garrard, A., 8
Chromatophors, 90
Clarke, A.C., 189
Cloak
 acoustic, 99–103
 earthquake, 103–105
 invisibility, 6, 8, 12, 15, 58–64, 67, 91, 94, 97–100, 102, 104, 115, 119, 123, 127–129, 131, 133, 140, 145–148, 153–155, 158, 161, 163, 168, 171–174, 176–180
 optical, 89, 90, 97, 98, 102, 131, 155, 158, 159, 162, 163, 170, 180
 water, 98, 103, 104
The Coming of the Fairies, 16
Copernicus, N., 37, 38, 45, 51
Cottingley Fairies, the, 16–17
Crab Nebula, 31
Cummer, S., 99, 101
Cuttlefish, 89–91

D

Dali, S., 10, 24
Dark matter, 82, 84

DARPA. *See* Defense Advanced Research Projects Agency (DARPA)
 Da Vinci, L., 26
 Dawkins, R., 27
 Defense Advanced Research Projects Agency (DARPA), 145, 146, 148
 Degaussing, 23
De Magnete, 107, 108
 Democritus of Abdera, 32
De Revolutionibus, 38
 Descartes, R., 38–40, 42, 43, 49, 50
 Deutsch, D., 18
 Dexler, E., 156
 Dielectric, 77, 93, 109, 119, 154, 158, 159, 166, 169
 Diffraction limit, 94, 148, 149
Dioptrica, 37, 39
Discourses Relating to Two New Sciences, 35
 Dodgson, C., 153
 Doppler effect, 139
 Doyle, A.C., 16, 17, 30

E

Eddington, A.S., 84
 Edgerton, H., 77
 Einstein, A., 81–85, 117, 126, 127, 184
 Einstein ring, 81–83
Electricity and magnetism, 51, 107, 111, 112, 122
The Emperor's New Clothes, 15
Encyclopedia Britannica, 20
 Ergin, T., 168
 Escher, M.C., 24
 Euclid of Alexandria, 32
 Event horizon, 12, 13, 160
 Everett, H., 17, 18
Experiments and Observations on Electricity, 111

F

Fang, N., 100
 Faraday, M., 111–116, 174
 Faraday's cage, 115
 Faraday's law, 113, 122
Fata Morgana, 73–74, 79
 Fermat, Pierre de, 40–42, 163, 181
 Fermat's principle, 40–42, 181
 Feynman, R., 156
 Fiber optics, 69, 86–90, 187
 Films
 Forbidden Planet, 3
 Invisible Man, the, 4, 5, 14, 20
 Silence of the Lambs, 3

The First Men in the Moon, 14
Flatland: a romance of many dimensions, 11
 Flying Dutchman, the, 73
 Foucault, L., 50, 51
 Fourier, J.B.J., 46
 Franklin, B., 109, 111, 115

G

Galilean moons, 43, 44
 Galilei, G., 35, 37, 43, 45
 Gamma Draconis, 45
 Gamma rays, 122–124, 127
 Gauss, C.F., 115–116
 Gauss's law, 115, 116
Genesis, 29, 77
Geometry and Light: The science of invisibility, 146
 Gilbert, W., 107–109
 Goldilocks effect, 17
 Grbic, A., 151
 Green, D., 196
 Green flash, 76–78
 Greenleaf, A., 179, 180

H

Harmonius Mundi, 37
Harry Potter and the Chamber of Secrets, 6
 Hawaiian bobtail squid, 90–91
 Hawking radiation, 11, 13, 187
 Hawking, S., 13, 14
 Head-mounted display (HMD), 67, 69, 71
 Hecht, E., 41
 Heiligenschein, 64
 Heinlein, R.A., 9
 Hero of Alexandria, 32
 Hertz, H., 122, 126
 HMD. *See* Head-mounted display (HMD)
 Holbein, H., 25
 Hooke, R., 36–37, 42, 48, 51, 54, 55, 80, 91
 Hsu, S., 13
 Huygens, C., 37, 38, 42, 48, 51, 118

I

Illustrated London News, 85, 113
 Impossibilities
 Type I, xii, xiii
 Type II, 9, 12, 13, 123, 127, 179
 Type III, 9, 13, 18, 127, 179
 Inami, M., 4

Interference

- constructive, 47–49
- destructive, 47–49, 133

Invisibility

- carpet cloak, 97, 99, 161–171, 174
- cloak, 6, 8, 12, 15, 58–64, 67, 91, 94, 97–100, 102, 104, 115, 119, 123, 127–129, 131, 133, 140, 145–148, 153–155, 158, 161, 163, 168, 171–174, 176–180

Invisible ink, 6–7

The Invisible Man, 4, 5, 14, 20

The Island of Doctor Moreau, 14

J

Journals

- Applied Physics Letters*, 153
- Cambridge and Dublin Mathematical Society*, 93
- International Comet Quarterly*, 196
- Journal of Optics*, 186
- Journal of the European Optical Society*, 95
- Nature*, 88, 145, 146, 156, 164, 166–168, 181
- Nature Communications*, 169, 170
- Nature Materials*, 91, 93, 146
- Nature Photonics*, 158, 167
- New Journal of Physics*, 94, 101, 102, 151, 160
- Optics Express*, 162
- Personal and Ubiquitous Computing*, 89
- Philosophical Transactions of the Royal Society*, 119
- Physical Review*, 20
- Physical Review Letters*, 97, 103, 141, 146, 171, 179, 181, 184
- Physics*, 176
- Science*, 101, 141, 146, 147, 163, 166, 168, 169, 175, 176, 178
- Science Express*, 153, 178
- Soviet Physics Uspekhi*, 134

K

- Kaku, M., 179
- Kepler, J., 37
- Kerr, J., 77
- Kitab al-Manazir*, 39
- Kusama, Y., 21, 22

L

Large Hadron Collider (LHC), 11–12

Law

- reflection, 32–33, 40, 63
- refraction, 32–33, 38–40
- Snell's, 38–40, 66, 136

Lens

- Eaton, 91–94
- gradient index (GRIN), 80, 91, 94
- gravitational, 81–83, 85
- hyper-, 149
- Luneburg, 94–96
- Maxwell's fish-eye, 93–94, 151, 152, 185
- super-, 93–94, 119, 120, 136, 140, 148–152, 168

Leonhardt, U., 92–94, 96, 136, 146–148, 151, 152, 175

Leucophores, 90, 91

Lewis, C.S., 15

Leyden jar, 89, 108–111

LHC. *See* Large Hadron Collider (LHC)

Light

- duct, 86
- propagation, 32, 42, 127, 129, 181, 182, 185
- ray, 32, 33, 40–42, 50, 51, 57, 64, 66, 79, 80, 87, 89, 92, 93, 95–98, 136, 137, 158, 159, 163, 166, 175, 176, 182, 186, 187
- speed, 11, 13, 20, 35, 38–40, 42–47, 50, 51, 78, 91, 121, 123, 125, 127, 134, 184

Linde, A., 18

Lockheed F-117, 132

Lodestone, 107–109

Lord of the Rings, 2, 3

Luneburg, R., 94–97

Lycurgus, 156–158

M

- Mach, E., xii
- Magritte, R., 19, 62, 72, 187
- Malus, E.L., 52, 53
- Malware, 19
- Mariotte, E., 28
- Marlowe, C., 5, 6
- Maxwell, J.C., 51, 52, 93–94, 107–129, 134, 135, 151, 152, 174–176, 178, 185
- Maxwell's equations, 123, 135, 174–176, 178
- McCall, M., 186, 187
- Meek, Sterner St. Paul, 3

Metamaterial, 91–93, 99, 100, 104, 119, 120, 123, 128, 131–187
 Michelson–Morley experiment, 123, 124, 126, 127
Micrographia, 36–37, 80
 Microwave, 29, 52, 93–95, 99, 115, 136, 141–144, 148, 151–155, 160, 163, 164, 169, 170, 177
 Mie, G., 8
 Mirage
 anamorphic, 84
 inferior, 78–80
 Novaya Zemlya, 86
 superior, 78–80
 Möbius strip, 182
 Moon illusion, 76
 Mortensen, N., 95
 Multiverse, 17–18, 181, 184
 Musschenbroek, P.v., 110

N

Nanobots, 156
 Nanowire, 151, 157–160, 164, 166, 167
 Narimanov, E., 158, 159, 181
 Newton, I., 36–38, 40, 42–43, 46, 47, 49–51, 117
 Nin, A., xi
 Noir, V., 65
 Nollet, A., 110

O

Oersted, H., 111, 112
The Once and Future King, 131
Opticks, 42–43
Opus Maius, 34
 Orosz, I., 25, 26

P

The Physics of the Impossible, xii
 Paintings
 The Ambassadors, 25
 The Crucifixion
 (*Corpus Hypercubicus*), 10
 The Human Condition, 62
 Mirror Anamorphosis with Column, 26
 Slave Market and the Disappearing Bust of Voltaire, 24
 The Treachery of Images, 19, 187
 Parallel universe, 17

Pascal, B., 58
 Pendry, J., 135, 136, 138–142, 146–148, 151, 158, 171, 175, 176, 178
 Penrose, R., 24, 135, 183
 Pepper's ghost, 59–60, 70
 Pepys, S., 36
 Periscope, 57
 Permeability (μ), 119–121, 134, 135, 141, 143, 176–178
 Permittivity (ϵ), 119–121, 134, 135, 141, 143, 176–178
 Philadelphia Experiment, 23
 Philbin, T., 146
 Photon, 91, 126, 168, 182
The Physics of the Impossible
 Planck, M., 126
 Planck scale, 12
 Platform 9 3/4, 146, 171
 Plato, 2, 3, 32
 Polarization, 52–54, 77, 97, 113, 158
 Pope Clement IV, 34
 Potter, H., 6, 146–148, 168, 171, 179–180
Principia Mathematica, 37
Profiles of the Future, 189

Q

Quantum gravity, 12, 13, 181

R

Radar, 30, 31, 47, 94, 95, 123, 131–133, 139, 140, 144–146, 170
 Radio waves, 122–124, 127, 171
 Rayleigh, L., 20
 Rayleigh scattering, 76
Red Riding Hood, 21, 22
 Refraction
 in air, 4, 33, 40, 66, 67, 73, 78, 163
 atmospheric, 75, 78, 79
 law of, 32–33, 38, 40
 Refractive index, negative, 120, 133–139, 141–146, 148–153, 163, 165, 166, 168, 171, 172
 Relativity
 general, 14, 81, 82, 84, 85, 127, 179, 181, 184
 special, 127, 181, 184
Republic, 2, 3
 Retroreflector, 64–68, 70, 89, 91
 Reutersvard, O., 23
Ring of Gyges, 2, 3

Ring of the Nibelungen, 2

Romer, O., 42–45

Rowling, J.K., 6

Royal Society, 36, 42, 46, 53, 55, 84,
117, 119

Rucker, R., 12

S

Schultheiss, V., 184

Schurig, D., 99, 101, 148, 153, 154, 175–178

Semouchkina, E., 153, 155

Shakespeare, W., 6

Shaw, P., 66

Siderius Nuncius, 43

Slade, H., 14, 15

Smart Shirt, 89

Smith, D., 14, 142, 146, 148, 153, 163, 164,
175, 176, 178

Smolyaninov, A., 181, 183, 184

Smolyaninov, I., 181, 183

Snell, W., 38

SONAR, 99, 100

Sonic crystals, 100–102

Spacetime, 127, 128, 179, 181, 184, 186

Spacetime cloak (STC), 186, 187

Split ring resonator (SRRs), 93, 141–145,
153, 157, 158, 163, 168, 177,
178, 181, 182

Stellar

aberration, 45

parallax, 45

String theory, 12

Strutt, J., 20

A Study in Scarlet, 30

Superconductor, 120

Sutherland, I., 67

Swift, J., xii

T

Tachi Labs, 60, 61, 67–72

Tachyons, 184

The Tempest, 6

Tesseract, 9, 10

The Time Machine, 14

Titanomachy, 2

Tolkein, J.R.R., 2, 3

Toulmin, S., 118

The Tragical History of Dr. Faustus, 6

Trait de la Lumiere, 48

Transcendental Physics, 14, 24

Transformation optics, 163, 169, 174–178

Treatise on Light, 38

Triboelectricity, 109

Trojan horse, 9, 19, 21, 23, 28, 71, 170

Twain, M., 89

Tyc, T., 92, 93, 96

U

Unruh, W., 191

U.S.S. Eldridge, 23

V

Vacuum pump, 53–56

Valentine, J., 164, 166, 167

van Musschenbroek, P., 110

Veer, Gerrit de, 86

Velocity

group, 138

phase, 138

Veselago, V., 134, 136, 139, 142

Virgil, 19

The Voyage of the Dawn Treader, 15

W

Ward, A., 136, 175

War of the Worlds, 14

Warp speed, 127–129

Watson, J., 30

Wave

amplitude, 46, 47, 49, 150

electromagnetic, 119, 122, 135, 138, 142,
143, 146, 148, 160

frequency, 46, 47, 123, 126, 138

packet, 138, 139

particle duality, 126

sound, 102

velocity, 134

wavelength, 8, 46, 47, 52, 76, 80, 94,
122–124, 133, 138, 139, 143, 145,
148, 149, 151–153, 155–160, 163,
166–169, 172

Web sites

arXiv.org, 162

physicsworld.com, 152

www.adsabs.harvard.com,

www.sciencemuseum.org.uk, 57

www.wikipedia, 65, 114, 124,
128, 134

Wells, H.G., 4, 14, 20

Whirlpool galaxy, 54
White hole, 13
Wired Magazine, 180
Wood, R., 20, 80
Wormhole (EWH), 179, 180
The Worst Journey in the World, 8

X

X-ray, 71, 131

Y

Young's slits, 48, 49
Young, T., 48, 126

Z

Zhang, B., 97, 171, 174
Zhang, X., 149, 151, 163, 165, 181
Zollner, J., 14, 16, 24
Zollner's illusion, 25