



Fifty Years of Relativity
Neutrons and Atom Patterns in Crystals
Nature Conservation and Ecology
Lighting and Work
The Automatic Factory
Research Report



25
AST 1, 10, 15
15E 2, 3, 4
16 2, 3, 4

SCIENCE NEWS

EDITED BY A. W. HASLETT

37

PENGUIN BOOKS

THE PRESS AND THE SCIENCE
OF THE FUTURE

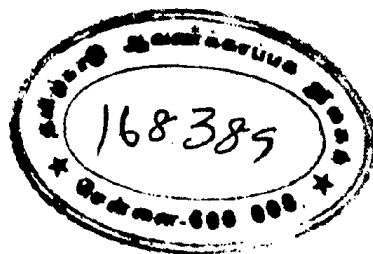
72
First published August 1955

A 1, N 46

N 55-37

Subscriptions to
SCIENCE NEWS
(which appears quarterly in
February, May, August, and November)
and

NEW BIOLOGY
(which appears half-yearly in
April and October)
are accepted at 11s each
for four issues
including postage



Made and printed in Great Britain
for Penguin Books Ltd
Harmondsworth, Middlesex
by The Campfield Press, St Albans

CONTENTS

FIFTY YEARS OF RELATIVITY W. B. Bonnor	7
NEUTRONS AND ATOM PATTERNS IN CRYSTALS Kathleen Lonsdale, F.R.S.	25
NATURE CONSERVATION AND ECOLOGY J. B. Cragg	43
LIGHTING AND WORK J. B. Collins	63
THE AUTOMATIC FACTORY H. D. Turner	90
RESEARCH REPORT A. W. Haslett	111
CORRESPONDENCE	122
SOME BOOKS RECEIVED	124
ABOUT OUR CONTRIBUTORS	127

26 illustrations are between pages 64 and 65

FIFTY YEARS OF RELATIVITY

W. B. BONNOR

RELATIVITY is now fifty years old. The first foundations of the subject were laid in a work by Einstein published in the autumn of 1905. Since then, the theory has brought about a revolution in physics and philosophy; and it has also been indirectly responsible, by changing our ideas of the nature of energy, for developments which can alter the whole course of civilization. Yet in spite of the fame it has achieved, and of the cogency which the specialist can discern behind the formidable mathematics, the theory has in some respects not fulfilled the expectations which it once aroused, and some physicists in recent years have thought that Einstein's relativity needs modification in certain ways.

In this article I shall trace the history of relativity since its inception fifty years ago, and describe some of its achievements and difficulties. At the outset I should explain that the theory falls naturally into two parts – the special theory, which has responded satisfactorily to nearly all tests, and which is almost universally accepted; and the general theory, for which there is less evidence, and which may still be regarded as controversial in some respects.

THE SPECIAL THEORY OF RELATIVITY

An important concept in nineteenth-century physics was that of the aether. This was supposed to be a fluid pervading all matter and space, but not resisting motion through it. It fulfilled two main functions. First, it was the medium for vibrations of electromagnetic waves, such as those of light. Secondly, it constituted a standard of reference for the absolute motion of bodies, since on the whole it was supposed to be at rest. One may picture the aether as an ocean whose surface is disturbed by waves, but which is otherwise still; the speed of a boat relative to various

observers in motion will be different, but all these observers can agree in assigning to the boat a definite speed relative to the water. If the aether is regarded as analogous to the water in the ocean, its functions can readily be understood.

An absolute standard for the measurement of motion was important in nineteenth-century physics for two reasons. In the first place, light was known to have a speed of 3×10^{10} centimetres per second, and it was supposed that this velocity must be considered relative to the medium in which the vibrations took place (i.e. to the aether), and not to the source of light. It was not easy to be sure of this because the speed of light is so great compared with terrestrial speeds, but from observations of light from double stars (i.e. pairs of stars which rotate about each other), it is now known that the speed of light is independent of that of its source.*

The second important need for an absolute standard of motion arose in the science of mechanics. According to Newton's laws of motion, every body remains at rest, or moves with constant velocity, unless it is acted upon by a force; and if a force acts on it, the acceleration produced is proportional to the strength of the force. These statements presuppose some absolute standards for measurement of velocity and acceleration. In fact, as Newton recognized, the laws of motion are true only for observers who themselves have no acceleration; for accelerated observers, they are false. The aether, it was thought in the nineteenth century, must be the absolute standard compared with which true velocities and accelerations are to be measured.

Towards the end of the century, many attempts were made to establish the existence of the aether. One possible way of doing this was to show that the earth had a velocity relative to it. The earth was known to be moving at a speed of 70,000 m.p.h. relative to the sun, so presumably it had also a speed relative to the aether, though, of course, not necessarily 70,000 m.p.h. because the sun itself might not be at rest in the aether. The most famous of the experiments to measure the absolute speed of the earth

*See the article by W. H. Marshall in *Science News* 34, pp. 61-2 especially.

was that of Michelson and Morley, carried out in 1887. These experiments (with one exception which I shall mention later) all gave the same result – that the earth had *no* velocity relative to the aether. Excluding the fantastic possibility that the earth was at absolute rest and the whole universe revolved around it, the inevitable conclusion was that the aether theory was false: the aether did not exist.

Since there was no aether, there was no absolute standard of motion. All velocities must be regarded as relative to given standards of reference, and of these standards none has any special claim to preference over any other. In particular, the motion of observers was relative: the absolute observers required by Newton's mechanics could not be defined, and Newton's theory was to this extent unsatisfactory.

This was the position at the beginning of the twentieth century. The first clear statement that this relativity of motion required a complete reformulation of the laws of physics came from the French mathematician H. Poincaré, who in 1904 enunciated the principle of relativity that *all the laws of physics must be expressed in such a way as to be valid for any observer, whatever his motion*. This principle, which is the logical consequence of the non-existence of absolute or privileged observers, is the fundamental postulate underlying both special and general relativity.

The relativity of motion was not easily accepted. For several years the aether hypothesis persisted in various forms. For example, to explain the negative result of the Michelson-Morley experiment it was suggested that the aether was dragged along with the earth so that experiments on the earth's surface would not be expected to yield the velocity relative to it. However, negative results from experiments performed on high mountains and in balloons have made this explanation untenable, even had it been acceptable on other grounds.

According to the principle of relativity the laws of physics must be the same for all observers. Now in Maxwell's theory of electromagnetism the velocity of light appears as a fundamental constant. If the equations of this theory are to be true for all

observers, the velocity of light must be the same for all observers (more strictly, for all observers not subject to forces: see below). This was what Einstein asserted in 1905; together with the principle of relativity it forms the basis of the special theory of relativity.

The special theory of relativity expresses the laws of physics in a form which is the same for a special class of observers, namely, those not subject to forces. Of course, in practice, observers always experience forces of one sort or another, but in many cases, even on earth, the forces are weak enough not to interfere very much with the application of special relativity. According to Newton's laws, forces cause accelerations, so the observers of special relativity correspond to what in the older language would have been called unaccelerated observers.

Between 1905 and 1910 much progress was made in building the new theory. As well as Poincaré and Einstein, the Dutch physicist H. A. Lorentz deserves much credit for the early developments, though even at his death in 1928 he did not accept the revolutionary results which followed from his work.

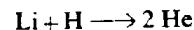
Many discoveries which have since become famous were made at this time. Among them was the famous relation between mass and energy which has become almost a commonplace since the explosion of the first atomic bomb. According to special relativity the energy E of a body is connected with its mass m (in appropriate units) by the formula

$$E = mc^2,$$

where c is the speed of light. This means that if energy is given to the body, for example, in the form of heat, its mass is increased. Since the number c^2 is so great, very large amounts of energy have to be imparted to the body to give it an appreciable increase in mass. For instance, if 2 kilogrammes of water are heated from freezing to boiling point, the increase in mass is only 10^{-8} gramme. For practical purposes, the relationship becomes important if considered the other way round: for a very small amount of the mass of a body will yield enormous quantities of energy, if the conversion can be carried out. The destructive power of the atomic bomb comes from the loss of about

0.1 per cent of the mass of atoms of uranium-235 originally present. This energy appears as light, and as kinetic energy of electrons, neutrons, and other fragments. If the fission is controlled and occurs in an absorber this energy can be trapped as heat; the new atomic power station in Cumberland is expected to produce in this way as much useful energy from 16 pounds of fissionable material as that obtained from 8,000 tons of coal. The energy for the explosion of a hydrogen bomb results from the slight decrease in mass which occurs when hydrogen is converted into helium, and it is the same process which is responsible for the outpouring of radiant energy from the sun and stars.

When the mass-energy relation was first put forward as a consequence of special relativity there was no direct evidence for it; but since then a great many experiments of nuclear physics involving transmutations of the elements with changes of mass have put the formula beyond reasonable doubt. As an example, we may consider the artificial disintegration of lithium (Li) with protons (H) into α -particles (He) according to the reaction



The masses concerned (taking as standard 16 for the oxygen atom) are

$$\text{Li} = 7.0182, \text{H} = 1.0081, \text{He} = 4.0039.$$

The decrease in mass in the disintegration is 0.0185 which is found to appear as kinetic energy of the α -particles.

The mass-energy relation was not the only spectacular result of special relativity. The theory required a fundamental change of outlook on space and time. It showed that two events simultaneous in one observer's experience would not usually be so to another who was in motion relative to the first; and motion was shown to affect even such apparently simple operations as measurement of lengths and times. For example, suppose there are two observers A and B equipped with identical clocks and rulers of length L centimetres. Now let B move in a straight line with velocity V relative to A ; then if B holds his ruler along his line of motion, it will measure according to A 's ruler, not L , but

$L\sqrt{1-V^2/c^2}$. The moving ruler appearing contracted, a phenomenon which is called the Fitzgerald contraction. Moreover, a lapse of time recorded on *B*'s clock as *T* seconds will be measured by *A* as $T/\sqrt{1-V^2/c^2}$, so that *B*'s clock appears to go slow; this is known as time-dilatation. Of course, owing to the relativity of motion, the situation is completely reversible, and, to *B*, *A*'s ruler seems contracted, and his clock seems to go slow.

It is clear from these formulae that, if *B* had velocity *c* relative to *A*, his ruler would appear to have no length and his clock would seem not to be going at all. In fact this could not happen because no material object can travel with a speed equal to, or greater than, that of light. This is because the energy required to produce such a speed would be infinite. I should emphasize that this is true only for *special* relativity; this is a common source of misunderstanding of the theory of relativity. To an observer on our rotating earth some distant stars have apparent velocities which are greater than *c*; but this involves no contradiction because the observer, who is using a frame of reference which is rotating, is subject to centrifugal force, and is not one of the observers to whom the special theory applies.

The Fitzgerald contraction and the time-dilatation become important, and indeed, measurable at all, only at high velocities. In recent years such high velocities have been realized by the acceleration of nuclear particles in the cyclotron and similar machines, and direct confirmation of the relativity formulae has become possible. Other confirmation has come from high-velocity particles produced by cosmic rays. Here, for instance, is an illustration of time-dilatation which comes from the study of mesons. A certain type of meson is formed from cosmic rays in the atmosphere about 10 kilometres above the earth's surface. It is known from experiment that such a meson has a life-time of about 3×10^{-6} seconds; that is to say, after existing for that time it changes into something else. To travel 10 kilometres in this time the meson would have to travel at about 10 times the speed of light, which is impossible according to special relativity. Yet these particles certainly do travel this distance in the time because they are known to reach the earth's surface! The explana-

tion is as follows: 3×10^{-6} seconds is the 'proper life-time' of the particle, or the life-time which would be recorded by a clock travelling with it. To an observer on earth, this moving clock would go slow, and so the life-time appears to him longer than the proper life-time. This explains why the particle seems to last longer than it should. A calculation shows that the speed of the meson is in fact about $0.99c$, and no contradiction arises. Alternatively, the paradox can be explained by the Fitzgerald contraction, because to a hypothetical observer moving with the meson the distance of 10 kilometres appears contracted by a factor of about 10, and can be covered in the proper life-time without exceeding the speed of light.

Evidence now leaves very little doubt about the correctness of special relativity. The only reliable experiments which seem to contradict the theory are a series of repetitions of the Michelson-Morley experiment by the American physicist D. C. Miller, who found a small but definite value for the velocity of the earth relative to the aether. These results have never been explained, although an attempt was made by J. L. Synge to reconcile them with special relativity by a detailed consideration of Miller's experimental set-up. Synge's work was not successful, but leaves the hope that a reconciliation along these lines might perhaps be possible. At the present time the possibility of repeating the Michelson-Morley experiment more accurately using radar waves instead of light is being discussed, and the results of this experiment should be of great interest.

GENERAL RELATIVITY

Special relativity made it possible to write physical laws in a form that applied to any observer not acted upon by forces. General relativity removed this restriction and showed how to express laws in a way completely independent of the motion of the observer. It thus overcame the difficulty in Newton's theory – that it applied only to certain (unaccelerated) observers.

To give some idea of the scope of general relativity, we must briefly consider the Newtonian theory which it replaced. This consisted of two parts. The first part is the dynamics, which

describes how bodies move when known forces are applied to them, and is contained essentially in the equation

$$P = mf, \quad (\text{Eqn 1})$$

where P is the force, which when acting on a body of mass m produces acceleration f . The second great contribution of Newton to mechanics was the law of gravitation, which states that the gravitational force F between two bodies is directly proportional to their masses m , M , and inversely proportional to the distance between them:

$$F = GmM/r^2, \quad (\text{Eqn 2})$$

where G is the universal constant of gravitation. In Newton's mechanics these two aspects, dynamics and gravitation, are really quite separate, but in Einstein's general relativity they are combined together into one theory.

To achieve this unification Einstein made use of an important fact which found no explanation in Newton's theory. Suppose that the force P in Equation 1, under which the body of mass m is moving, is the gravitational force F , given by Equation 2. This will be so, for example, for the motion of the earth (of mass m) under the sun's gravitational attraction, if M represents the mass of the sun. We then write

$$mf = GmM/r^2,$$

and, cancelling the ms , we have

$$f = GM/r^2.$$

This gives the earth's acceleration, and from it we can calculate the earth's orbit. But are we really justified in cancelling the ms ? The m in Equation 1 represents the inertia of the body concerned, that is, its resistance to a force tending to accelerate it. This m is called the *inertial mass* of the body. The m in Equation 2, however, has nothing to do with inertia, but measures the power of the body to cause a gravitational force, or to respond to the gravitation of other bodies; it is called the *gravitational mass* of the body. If you start to push a motor car on a level road, your effort is required to overcome the inertial mass, and has nothing to do with gravitation; but if you lift a heavy weight, it is the gravitational mass you have to cope with. The two masses really represent quite different concepts.

It is known from experiment that the gravitational mass of a body is always equal to its inertial mass; indeed this is a consequence of the fact, known to Galileo, that all bodies, whatever their size, would fall to the earth in a vacuum at the same rate. But Newton's theory does not explain why this is so. With characteristic genius, Einstein took this experimental result, and used it to unify dynamics and gravitation.

General relativity is founded on two hypotheses: the principle of relativity, and the equality of gravitational and inertial mass. The mathematical formulation of the theory is very different from that of Newton, and, in fact, of any previous theory. This formulation is essentially a geometrical one and is an extension of the geometrical approach to special relativity introduced by Minkowski who showed that the special theory was much simplified if the three dimensions of space and one of time were treated together in a single four-dimensional continuum. With the abandonment of absolute simultaneity it had become clear that Newtonian ideas of absolute space and time had to be given up, and in certain sense, space and time formed a unified whole. The four-dimensional space-time of Minkowski was the natural expression of this.

Just as the relations between the points of three-dimensional space are studied by geometry, so may be relations between points in space of four (or any number) of dimensions. Usually such geometries have no physical application, but nevertheless they are interesting to pure mathematicians, who, unlike physicists, feel no obligation to confine their studies to the external world. In fact, the sort of geometry needed by general relativity had been invented as pure mathematics by Riemann long before any physical use was contemplated for it. Einstein was able to use Riemann's work to express mechanics in geometrical form.

According to general relativity not only are accelerations relative, but so also is gravitational force. Indeed the notion of gravitational force plays no part in general relativity. (I shall later refer to attempts to eliminate force completely.) The fundamental entity is the space-time continuum. The problem of astronomy, in which bodies move under their mutual gravita-

tional attractions, becomes one of the geometry of space-time. If the nature of this geometry is known, this by itself determines the motion of the stars and planets, and there is no need to consider the gravitational forces between them.

Einstein put forward the general theory of relativity in 1915, and its importance was recognized at once. It was found that for a wide range of phenomena the new theory gave substantially the same results as Newton's theory; as the latter was known to be very largely correct, this was itself highly satisfactory. There were, however, three astronomical predictions in which relativity differed from the older mechanics, and these came to be known as the 'three crucial tests of general relativity'.

The first crucial test concerned the orbit of the planet Mercury, which had been known for many years not to agree with Newtonian theory. Calculations soon showed that the field equations of general relativity gave precisely the correct orbit. It has since been found that the orbit of the earth shows a similar, but smaller divergence from the Newtonian prediction, and here too relativity gives the correct result.

The second test is called the bending of light. According to general relativity, rays of light are bent when they pass through a strong gravitational field. This is really not surprising when one remembers that according to the mass-energy relationship, light, having energy, also has mass, and therefore responds to (and exerts) gravitation. If a ray of light from a star passes near the sun it is deflected from its straight-line path, and to an observer on earth the star appears slightly displaced from its true position. Normally the sun's own light is too strong for us to see stars in this position, but they become visible during a total eclipse. The true position of the star can be obtained by observing the same part of the sky at night some months later when the sun is no longer between the star and the earth. By a comparison of the true and apparent positions of the star the amount of bending can be measured. The first measurements were made in 1919 and seemed to agree with the relativity prediction. Further apparent confirmation came from the eclipse of 1922, but measurements during eclipses in 1929 and 1947, and also a

re-examination of the earlier results, suggested that the relativity value is about 10 or 20 per cent too small. On the other hand, the observations of 1952 agreed well with relativity. The bending of light is very difficult to measure accurately, and statistical methods have to be used to get the most probable result from a set of observations. Most elaborate precautions were taken by an expedition led by Professor E. S. Freundlich of the University of St Andrews Observatory to get accurate measurements of the 1954 eclipse, but unfortunately no observations at all could be taken because of the weather. The present position about relativity and the bending of light seems to be that the observations indicate a possible discrepancy, but they are certainly not reliable enough to disprove the theory. I should add that the Newtonian value for the bending of light (half the relativity value) is certainly wrong.

The third crucial test is called the gravitational shift of the spectral lines. General relativity predicts that the wavelength of light emitted by atoms, such as those on the sun, which are in a strong gravitational field, will be longer than the corresponding wavelength on earth. At first the observations from the sun and also from a very dense star, called the companion of Sirius, seemed to confirm the relativity prediction, but more recent observations by Freundlich and other astronomers on these and other stars have shown that the evidence does not conform with relativity. It seems probable that the gravitational shift is a complicated phenomenon, depending on other factors besides gravitation, and at present no evidence for or against relativity can be adduced from it.

Further evidence bearing on general relativity came later from a different source. In 1917 Einstein proved that if an additional term (the cosmological term) were added to his field equations a solution of them could be found which referred to the universe as a whole, or at least, to a simplified model of it; and in the same year W. de Sitter found another solution corresponding to a different world-model. This was the beginning of the science of cosmology, which has turned out to be one of the most fascinating studies of modern times. I will give a very brief

Einstein and many others have believed that it must be possible to produce a theory in which gravitational, electromagnetic, and nuclear forces are more closely unified than this, and in which all three types of force are expressed in the structure of the space-time continuum. The search for this – a unified field theory as it is called – absorbed much of Einstein's energies during the last thirty years of his life. He was guided throughout by a belief that the attitude of quantum theory – that in atomic physics only probabilities can be predicted – is fundamentally wrong, and that the correct physics must be deterministic. As he put it: 'I cannot believe that God plays dice with the world.'

Yet all the efforts to create a unified field theory have so far been unsuccessful. The most recent attempt, developed independently by Einstein and E. Schrödinger during the last ten years, has in many ways come nearest to success, but in its present form it is not a satisfactory solution of the problem. The failure to find a unified field theory in spite of so much effort leads one to doubt whether the geometrization of physics can be carried through at present. Now that the more obvious extensions of general relativity have been tried, it seems probable that even if a satisfactory unified theory could be formulated it would be too complicated to be workable.

OTHER RELATIVITY THEORIES

The lack of unified field theory, though disappointing, is not a valid objection to general relativity. This may still be considered in its own right as a dynamics and gravitation theory, and is to be judged by its success in these fields. I shall now refer to some criticisms of general relativity which have been made, and to some of the alternative theories which have been proposed.

One of the difficulties in general relativity concerns the method of making measurements. We have seen how in special relativity rulers and clocks are affected by motion: in general relativity measuring instruments are affected also by gravitational fields and accelerations. For example, Professor W. H. McCrea has shown that an observer who is measuring lengths

in a gravitational field may get different results from his ruler if he changes his hold on it from one end to the other. To define what is meant by a rigid ruler becomes a difficult matter, and the motion of distance is accordingly no longer a simple one.

These relativistic effects are ordinarily so small that these problems concerning measurement involve no practical difficulties, and it may be possible to overcome them without modifying general relativity. However, they have led some workers, notably E. A. Milne, and the philosopher A. N. Whitehead, to construct different theories in which measurement is a more straightforward process. Whitehead's theory is formulated only for a certain class of observers so that it does not fully implement the principle of relativity. It represents a sort of compromise between the theories of Newton and Einstein, and it gives the same results as general relativity for the three crucial tests. It suffers from the disadvantage that the observers to which it applies are not properly defined; like Newton's theory it requires its observers to be unaccelerated without explaining how an absolute acceleration is to be specified.

Milne's theory, first put forward in 1933, is a much more radical departure from general relativity. Milne asked the question: What is the simplest and most logical way for observers to carry out measurements of space and time? He came to the conclusion that, given clocks (which need not be the same for them all) and theodolites (for measuring angles) the observer could do without rigid rulers for measuring distance. An observer could measure distance by the radar method: he could send out light signals and obtain the distance of a given object from the time the signal took to be reflected and to return to him. In fact, Milne *defined* distance as the product of this time multiplied by the velocity of light, taken as a fundamental constant of nature. Milne used this definition of distance, certain primitive ideas about time, and the cosmological principle, as bases for a deductive theory, known as kinematic relativity, in which many of the results of ordinary relativity appear merely as consequences of the method of measurement.

Valuable though it is as an epistemological study, Milne's

theory is unsatisfactory as a substitute for general relativity. It is open to the obvious criticism that measurements of distances are not, in practice, made by means of reflected light signals, and it does not give a satisfactory gravitation theory; indeed it does not give the right orbit of Mercury. Its value has been to show how important the study of measurement can be, and to suggest that the laws of the universe may be more dependent on our ways of observing it than we usually think.

In Milne's theory, like Whitehead's, laws of nature are formulated for a special class of observers; but in kinematic relativity, this class is properly defined, and consists, in fact, of the cosmic observers referred to earlier. Thus to Milne cosmology and the cosmological principle are fundamental, and appear at an early stage in the development of his theory.

H. Bondi, T. Gold, and F. Hoyle, who advocate what is called the steady-state model of the universe, also attach great importance to the cosmological principle. Their argument is that it is only this principle which permits us to extrapolate our local laws of nature to remote regions of the universe; and they adopt a more powerful version of it, called the perfect cosmological principle (p.c.p.) which asserts that the universe presents the same large scale aspect to a cosmic observer at all places and at all times. This means that the average picture of the universe is unchanging in space and in time. In the world-models of general relativity the expansion of the universe means that the average density of matter decreases with time. This contradicts the p.c.p., and in the steady-state model, matter is supposed to be continually created out of nothing to keep the density constant.

The continuous creation of matter violates the principle of conservation of energy, but Bondi, Gold, and Hoyle are prepared to accept this for the sake of a model which complies with the p.c.p. A mathematical expression of the theory has been given by Hoyle, who alters the field equations of general relativity so that continuous creation is allowed.

The steady-state theory was designed to overcome a difficulty in the world-models of general relativity which no longer exists. Until a few years ago it was thought that these gave a value for

the age of the universe (T of the previous section) which was too small to satisfy other astrophysical evidence. The latest observations show that this not so. But, although its original *raison d'être* has disappeared, the steady state theory is still interesting for its attempt to come to grips with the problem of creation of matter, which in other theories is removed from consideration because all matter is supposed to have been created at once at the time of infinite density, T years ago. It is also possible that the theory may be able to explain the formation of galaxies, which in general relativity is at the moment an unsolved problem.

As an alternative to general relativity, Hoyle's mathematical formulation of the steady-state theory has the disadvantage that it seriously complicates Einstein's field equations. This is a result of abandoning conservation of energy, which is, as it were, built into Einstein's equations, and is one of their most beautiful and powerful features. It seems, at least to me, that a great deal more astrophysical evidence will have to be forthcoming before Hoyle's equations will be acceptable in preference to Einstein's. However, as a result of work by McCrea it may be possible to reconcile the steady-state theory with relativity. Indeed, the steady-state model does satisfy Einstein's field equations, but only if one allows the existence of a negative pressure in internebular space. The pressure of a gas or liquid is normally positive, but it is just possible that the negative pressure required in McCrea's work may find a reasonable interpretation in some quantum process associated with the continuous creation which in the steady-state theory is supposed to be going on in internebular space. (For a further account of the steady-state theory, see the article by H. Spencer Jones in *Science News* 32.)

CONCLUSION

The history of science suggests that no theory of applied mathematics or physics is likely to be permanently 'true'. In this respect these subjects differ from pure mathematics, which may be, in some sense, absolutely true. The work of the applied

mathematician is to be judged rather by its influence on the stream of scientific and philosophic thought. The influence of relativity has been very great, and, it seems safe to say, beneficial in that it has given impetus to other branches of thought, and in that its ideas are likely to be of lasting value. From this point of view Einstein's relativity would still have been one of the greatest theories of physics, even if it were proved wrong tomorrow. How much longer it will in fact be before the theory is superseded, nobody can, of course, be sure. But my own guess is that Einstein's relativity will be with us for some time to come.

NEUTRONS AND ATOM PATTERNS IN CRYSTALS

KATHLEEN LONSDALE, F.R.S.

WHENEVER the scientist develops a new technique, it is fascinating to speculate as to what new knowledge he will gain by using it. The results at first may be slow in coming and of interest only to the specialist. But the time soon comes when it pays scientists in other fields to sit up and take notice. Review papers are then given at a meeting of one of the 'Associations for the Advancement of Science' or are published in one of the scientific journals having a wide coverage of scientific news. Such papers may well be written, as this article is, by someone not actually using the new method, but qualified to assess it critically and to judge the importance of the information gained, from the point of view of some other allied scientific fields.

The use of neutrons to study the arrangements of atoms in solids was begun on a serious scale in 1946, when a good deal was already known both about neutrons and about crystal structure.

This was a very different state of affairs from that existing in 1912 when the first experiments were carried out on the diffraction of X-rays by crystals. Then the nature of X-rays was not certainly known; it was the subject of scientific controversy. Some physicists said that X-rays were particles, others that they were electromagnetic waves, but of very short wave-length, about 1 ångström (1 Å). W. H. Bragg finally suggested that they must 'partake of the nature of both', and this assumption of a state of quantum co-existence finally proved to be correct.

The arrangements of atoms in solids were not known either. They were believed to form repeating patterns, and some shrewd guesses had been made about the nature of these patterns for

rock-salt and some other simple compounds, but some of these guesses were right and others were wrong.

The early experiments gave vital information not only about the nature of X-rays and about the atomic patterns in solids (whether in the form of single crystals, or in the form of powders or of massive materials composed of many disorientated crystal-lites), but also about the nature of the interaction between them. The electrons of the individual atoms scattered (redirected) the X-ray waves in all directions, but with a complete reversal of phase. Since the atoms were regularly arranged, the scattered waves reinforced one another in certain directions and destructively interfered in others, giving a diffraction pattern which could be recorded on a photographic film, or detected by means of the ionization produced by the diffracted X-ray beams. W. L. Bragg proved that the positions of the diffracted beams were given by the equation $n\lambda = 2d \sin\theta$, where n was an integer, λ the wave-length of the X-rays, d the distance apart of successive planes of atoms in any particular direction in the crystal, and θ the glancing angle of incidence on those atomic planes, at which reinforcement of scattered waves would occur. Figure 1 shows the historical proof of this formula. At the same time, since the electron distribution in the atom itself was spread over a certain volume, there was some destructive interference in the scattering even from a single atom, and this was most effective for large values of θ and for small λ , leading to a marked falling-off of intensity in the high-angle diffracted beams. This 'falling-off' became known as f_0 , the atomic scattering factor. It can be thought of as representing the difference in scattering behaviour between free electrons, which move independently, and those bound in atoms, whose scattering is always inter-related. The temperature vibrations of the atoms accentuated this falling-off, so that in general (and especially for short wave-length X-rays) little or no diffraction occurred at angles approaching 180° .

Since the X-ray scattering was by the atomic electrons it was possible to predict what the intensity of scattering by any given kind of atom would be; in fact it would be Zf times the scattering by a single electron, where Z is the atomic number, and f is

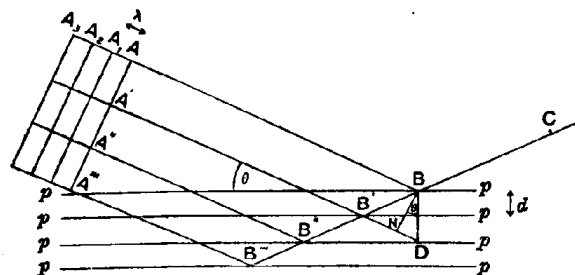


Fig. 1. Historical proof of the formula $n\lambda = 2d \sin\theta$. (W. H. & W. L. Bragg, *X-rays and Crystal Structure*. Bell, 1915.)

the atomic scattering factor which here includes the thermal vibration effects mentioned above, c the velocity of light. Hydrogen atoms ($Z=1$) are therefore least efficient, and heavy atoms such as Pb ($Z=82$) are most efficient scatterers of X-rays.

The absorption of X-rays by solids composed of various kinds of atoms was also studied and found to obey a relatively simple law, being rather large in general, but especially large for predictable wave-length and atomic conditions.

Finally it became possible to use the measured values of θ and λ to determine the geometry of unknown crystalline structures and to use the measured intensities of diffracted beams to determine the electron density distribution (and hence the positions of the atoms) throughout even very complicated solids. There was, however, a difficulty – that, though the relative intensities of different beams could indeed be measured, there was, in general, no way of telling whether the different beams were in step, or phase, with each other or not. Mathematical procedures of varying complexities were called in to assist in interpreting the X-ray patterns, and computing devices such as Hollerith machines were used to lighten the burden of arithmetical calculation involved. After a period of 40 years many thousands of atomic patterns were known with varying degrees of certainty.*

*For a further account of methods, see, for example, 'Determining Atomic Patterns' in *Science News* 18.

But there were certain gaps in our knowledge.

X-rays could give little information about the positions of hydrogen atoms. They could not distinguish between different isotopes (for example, those of hydrogen, H and D) since these have the same atomic number; and even when compounds of pure isotopes were made (ordinary and heavy ice, say) the differences found by X-rays were only second-order dimensional effects. Neither could they distinguish, in general, between atoms of closely neighbouring atomic numbers such as iron ($Z=26$) and cobalt ($Z=27$). And finally, there were certain physical properties of great practical importance to which X-ray diffraction appeared to be almost completely insensitive.

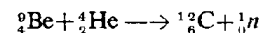
One group of such properties depends on some form of magnetic ordering of atoms. It is quite a common property of atoms that the individual atom should act as a miniature magnet. Several less common, but very interesting, properties of materials depend on the ability of particular atoms to align themselves magnetically, either parallel (and in the same direction) or antiparallel, within limited volumes of a crystal. In such cases there is thus a magnetic as well as a positional ordering of the atoms. The alignment is parallel in ferromagnetic materials, which offer the best known example. It is antiparallel in other cases, and the effects of the individual atomic magnets are then mutually cancelling, either wholly or in part. Both situations are of interest in relation to the theory of magnetic materials,* and the case of partial cancellation is also of technical importance. It offers the possibility of controlled magnetic properties.

Some years after the earliest X-ray experiments on crystals it was realized that, just as waves have particle-like properties in some respects (for example, in having their energy absorbed in finite amounts called quanta), so particles must have some wave-like properties. It was predicted, and triumphantly verified, that electrons could be diffracted by crystals. In this case, however, the scattering was not by the electrons of the atom, but by the whole field of force of the nucleus and electrons combined.

*See W. D. Corner's article 'Ferromagnetism' in *Science News* 30 (1953).

Moreover, the absorption of an electron beam was even higher than that of an X-ray beam, so that electron diffraction experiments were useful, in general, only for gases or for surfaces or thin films, although recently their use has been extended in special cases to investigations of crystal structure using methods similar to those employed in X-ray structural research.

When the neutron was discovered by Chadwick in 1932, it followed naturally that the possibility of its being diffracted by crystals would be considered quite soon, at least in theory; not at first as a practical technique, for the simple reason that it was not at first possible to get enough neutrons to do more than try the experiment qualitatively. They had to be knocked out of nuclei, one by one, so to speak, or obtained by the relatively inefficient interaction of helium nuclei (from radium) with beryllium nuclei



and it was not until the possibility of a chain reaction had been realized and atomic piles had been built that a sufficiently intense neutron beam could be obtained for quantitatively reliable experiments to be attempted; though the fact of neutron diffraction by crystals had been demonstrated by H. Halban and P. Preiswerk and by D. P. Mitchell and P. N. Powers in 1936.

By the time that atomic piles were available, a good deal was known about neutrons, one very important fact being that they are a dangerous form of radiation of great penetrating power. The technique of neutron diffraction by crystals therefore involves precautions far beyond those necessary when X-ray methods are used: heavy shielding, automatic controls, and very long counter chambers to collect all the diffracted neutrons. The precise nature of the neutron-crystal interaction was not known, but since many crystal structures of simple elements or compounds were known from X-ray studies, it was not difficult to obtain, empirically, the intensities with which neutrons are scattered by different kinds of atoms. This indeed was necessary, for although the Bragg law, $n\lambda = 2d\sin\theta$, holds equally well for neutrons and for X-rays, the relationship between neutron scattering and atomic structure has turned out to be rather

complex. This is partly due to the nature of neutrons themselves.

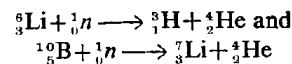
Neutrons are uncharged particles of size about equal to that of the hydrogen nucleus (the proton, H^+), and very slightly heavier. Although uncharged, they do have a magnetic moment, the property of a particle which determines the torque acting on it when in a magnetic field. Like the proton, too, they have an inherent 'spin' or angular momentum. More about these two properties of nuclear particles and their measurement can be learnt from the article by A. D. Little in *Science News* 25. The magnetic moment of the neutron is smaller and of opposite sign to that of the proton (-1.91 , compared with $+2.79$, nuclear Bohr magnetons); its spin is the same ($\frac{1}{2}$ in each case). As will appear later, both these properties may affect the interaction between neutrons and atoms that they encounter. The magnetic property becomes of special importance when the atoms encountered have themselves a resultant magnetic moment – and still more so when magnetic ordering occurs, a situation in which X-ray methods, as has been mentioned, have little to tell us.

Since neutrons are 10,000 times smaller than the atom itself and are uncharged, they can usually shoot straight through the atom without being scattered, unless the atom itself has a resultant magnetic moment. This case of magnetic interaction will be considered later.

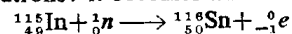
Another property of neutrons that is relevant is their velocity in the beam used. The equivalent wave-length of the neutron, as of other particles, depends on the product of mass and velocity, and the greater this product, the shorter is the wave-length. When first produced by fission, the neutrons have velocities of over 135,000 metres per second, with equivalent wave-lengths of $0.03 \text{ \AA}$ or less. These are too short for diffraction purposes and the effect on atoms of direct collisions with neutrons of such high energies would be a disruptive one – which is not what is wanted. In practice, these high-energy neutrons are slowed down by means of a succession of collisions within a moderator, such as heavy water or graphite, to thermal velocities (of the order of 2,000 metres per second) before being used for diffraction experiments. This means that they are in equilibrium with matter

at room temperatures or thereabout, and from the point of view of diffraction are equivalent to waves of length 1 to 2 $\text{\AA}$, just the same range as that common in X-ray diffraction studies, which is also approximately the distance between neighbouring atoms in crystals.

For the diffraction of neutrons to be of value, means are needed to detect the directions and to determine the intensities of the scattered neutron beams. With X-rays there is no problem; either a photographic film or a Geiger counter is a sufficient instrument of detection. Neutrons cannot be directly recorded photographically because they have no chemical action on a photographic emulsion, and their ionizing action is different from that of X-rays. Even thermal neutrons are not absorbed by most atoms to any extent. The light atoms of lithium-6 and boron-10 do absorb neutrons appreciably, according to the following relations:



and can therefore be used as detectors. In fact counters filled with boron trifluoride are the usual means of observing the positions and intensities of diffracted neutron beams. Other relatively strong neutron absorbers are cadmium, which is used as a shutter or screen, certain rare-earth atoms (especially samarium, europium, and gadolinium) and indium (In), which is made radioactive by neutrons: it becomes tin and emits a β -particle.



The β -activity induced in indium by the absorption of neutrons can be recorded photographically and in this way neutron diffraction patterns may be obtained on photographic plates or films; but in general the ionization method is that normally employed.

Even when there is no appreciable absorption, however, thermal neutrons which come close enough to the nuclei of other atoms will be scattered, that is, redirected from their original paths, and this scattering is of two kinds. In the first place neutrons, though practically unaffected by the much lighter electrons, are scattered by the nucleus just as if both were impene-

trable spheres, and yet in many ways just as if neutrons were also waves. This 'potential scattering', which like X-ray scattering is out of phase with the incident neutron beam, can be calculated from wave-mechanical theory and is expected to vary as r^2 , where r is the nuclear radius which is proportional to the cube root of the mass of the atom. This would imply a slow but steady increase in the intensity of neutron scattering with increase in the atomic weight of the scatterer. In fact, the measured values of the neutron scattering from the different atoms, though varying within a much narrower range than those for X-rays, are very irregular even for different isotopes of the same atom. Thus the amplitudes of neutron scattering (measured in units of 10^{-12} cm.) for ^{58}Ni , ^{60}Ni , and ^{62}Ni respectively, are 1.44, 0.30, -0.87 . (The convention is that scattered waves which are out of phase with the incident waves are taken as having positive amplitudes, those in phase, which is much more unusual, as having negative amplitudes.)

This irregularity, which may lead even to a negative amplitude of scattered neutron radiation, is due to the second kind of scattering. On impact the neutron and nucleus temporarily combine to form an unstable compound nucleus and then separate. This gives rise to a 'resonance scattering' which depends upon the energy levels of the compound nucleus and which may be greater or less than, and in phase or out of phase with, the potential scattering. The measured neutron scattering amplitudes are the sum (positive or negative) of both these effects. The resonance scattering due to the formation of the compound nucleus cannot be predicted, because we know too little about nuclear forces. Any condition which gives rise to different types of compound nucleus will give rise to different neutron scattering effects.

Since the neutron and the scattering nucleus are both extremely small (10,000 times smaller than the electron cloud around the nucleus) there is no interference within the atom: no atomic scattering factor, as in the scattering of X-rays by the electron cloud, although there is a falling off of intensity at high scattering angles due to thermal vibration.

The nuclear scattering of common and heavy hydrogen (H and D) are comparable in intensity with that from other atoms, and therefore much of the crystal structure investigation carried out by neutron diffraction has been directed towards that part of the subject which is most difficult by X-ray methods, the location of H atoms in structures containing heavier elements.

The matter has been complicated by the fact that the H nucleus, like the neutron, has spin $\frac{1}{2}$ and that in the temporary compound nucleus these spins may be aligned parallel or antiparallel, so that even H has two different scattering amplitudes. In a crystal, however, there is no spatial regularity, no periodicity, in the way such different compound nuclei will be formed (any more than in the arrangement of isotopes of any given element) and therefore the effect of this dual possibility is to introduce a degree of uncertainty, of 'incoherence', which diminishes the intensity in the diffracted beams and increases the general background intensity of the pattern as a whole.

This background intensity may be increased in various other ways as well as by spin incoherence: by the presence of more than one isotope, by multiple scattering, by thermal vibrations and by any other kind of irregularity in the structure such as a statistical distribution of different kinds of atoms (or of vacancies) in equivalent atomic sites, or of one kind of atom in different atomic sites. Thus a study of the background scattering may sometimes yield information as important as that given by the positions and intensities of the diffracted beams.

In another respect neutron diffraction studies are complementary to X-ray investigations. It often happens that atoms of close atomic numbers have quite different neutron scattering amplitudes (chromium, 0.35; manganese, -0.37 ; iron, 0.96; cobalt, 0.28; nickel, 1.03) and therefore the replacement of one such atom by another (in minerals such as the spinels) or the relative degree of order or disorder in an alloy of two or more of them, consequent on heat-treatment, ageing, or mechanical treatment, can be detected by neutrons, though not by X-rays.

The mention of these particular atoms brings us to the use of neutrons in the study of magnetic materials. As has been men-

tioned earlier, certain atoms behave as magnets, and such atoms can form several different kinds of magnetic crystal. The usual origin of atomic magnetism is that the atoms concerned have a number of electrons of certain energies which are not balanced (paired) by other electrons of equal energies and opposite spin, but there may also be some magnetic effects due to unbalanced orbital moments. In the simplest case, the atomic magnets are orientated in all directions at random. This invariably occurs at high enough temperatures, owing to thermal vibrations. The crystals are then said to be 'paramagnetic'. They are attracted weakly by a magnet, but do not acquire any remanent magnetism of their own.

If, however, the radial distribution of the unpaired electrons is related to the distance between neighbouring atoms within certain limits, then the atomic magnets line themselves up with their residual spins parallel. This condition is satisfied in the elements iron, cobalt, and nickel, as well as in certain rare-earth elements and some alloys. The atomic magnets are aligned only within certain volumes called 'domains', and the domains themselves within a single crystal may be randomly orientated in the absence of any external magnetic field. If an external magnetic field is imposed, however, all the spins in every domain will tend to reorient themselves so as to align the atomic magnets along the external field and will do so completely (to saturation) if the field is large enough. If the external field is removed they tend to remain aligned, although not with complete success. This whole phenomenon is known as 'ferromagnetism'. As the temperature is raised the ferromagnetism disappears, giving place to paramagnetism, at a definite temperature for each such substance: the Curie temperature.

Instead of the atomic spins being lined up parallel (with all N poles in the same direction), they may be lined up antiparallel (alternate N and S poles in the same direction) in domains, and in this case the crystal does not respond to an external magnetic field because the individual domains and, indeed, the individual unit cells of the crystal have no resultant magnetic moment. Yet there is a high degree of regularity in the spin orientations

and there is again a definite temperature, the Néel temperature, at which this 'anti-ferromagnetism' disappears, leaving only paramagnetism. Many oxides, sulphides, selenides, fluorides of ferromagnetic elements iron, cobalt, and nickel, and also of elements such as manganese and chromium – which do not quite satisfy the conditions for being ferromagnetic – possess this anti-ferromagnetism at low temperatures. Such crystals are formed of positive, metallic ions and negative ions of oxygen, fluorine, etc., and the negatively charged partner is believed to play an essential part in lining up the spins of suitable neighbouring metallic ions. The effect of the negative ion is greatest when it is centrally placed with respect to two metallic ions, to which it is united by valency bonds reaching out from it in opposite directions.

Finally there are compounds, of which the mineral magnetite, Fe_3O_4 , is the most famous example, in which some spins are aligned one way, some the opposite way, but the numbers are unequal, so that they do not balance out and there is a resultant ferromagnetism. Such substances are now called 'ferrimagnetics'. Many such crystalline compounds contain two or more different kinds of metal atoms, of which some are tetrahedrally surrounded by oxygen atoms and others (twice as many) are octahedrally surrounded by oxygen atoms, the spins being oppositely aligned in tetrahedral and octahedral sites respectively). By substituting one kind of atom for another, the magnetic properties of the material can be varied in a controlled way. Figure 2 summarizes diagrammatically the distinction between these three kinds of magnetic ordering.

Now X-ray methods can only detect changes in the magnetic order in a crystal as a second-order effect. These changes do often give rise to small distortions of the whole crystal lattice, and these small variations in dimension can be detected, and may even give some hint as to what is happening to the atomic spins.

For neutron scattering, however, the effect is a first-order one. Since the neutron has a magnetic moment of its own, the direction of the residual magnetic moment of the atom has a large

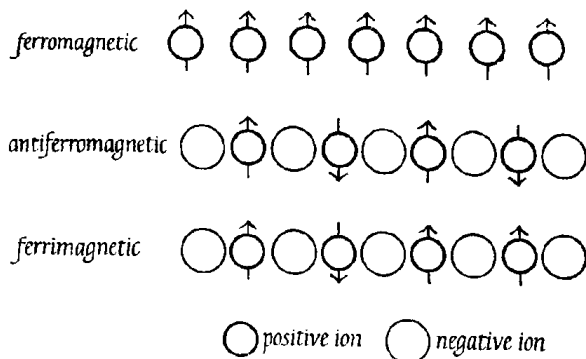


Fig. 2. Diagrammatic representation of ordered arrangements of magnetic ions.

effect upon the neutron scattering. For magnetic atoms, in addition to the rather complicated nuclear scattering, there will be a magnetic scattering of neutrons by the unpaired electrons. In paramagnetics, this is disordered and will contribute only to the incoherent background; in ferro-, antiferro-, or ferri-magnetics the alignment of atomic magnets is ordered and periodic, and diffraction effects are observed which, as it happens, are of the same order of magnitude as the nuclear scattering itself, but which can be distinguished from the latter without much difficulty.

In the first place, the magnetic scattering to be expected for any given angular relationship of atomic and neutron spins can be predicted. There is an atomic scattering factor similar to that found for X-rays, except that in the neutron case not all the electrons but only those with unpaired spins are scattering and mutually interfering within the atom itself. There is a variation with temperature, together with an abrupt disappearance of coherent diffraction when the crystal becomes paramagnetic. And in the case of ferro- and ferri-magnetism there is a variation of intensity of neutron diffraction with the direction of an applied external magnetic field relative to the orientation of the diffracting crystal planes of atoms.

The neutron diffraction studies have, in fact, been the principal method by which the directions of alignment of atomic moments in these magnetic materials have been studied and various theories tested or confirmed. In the case of some antiferromagnetics, such as manganous oxide, MnO (Figure 3), it has been found that the 'magnetic unit cell' found by neutrons has double the dimensions, in all directions, of the 'chemical unit cell' found by X-ray methods and in other cases, such as manganous fluoride, MnF_2 (Figure 4), the unit cell, the unit of pattern which repeated indefinitely gives the whole crystal, is the same for both. It is not necessary for the directions of all the atomic moments to be aligned, positively or negatively, in one direction only. In MnO and MnF_2 , this does appear to be the case, but in manganese dioxide, MnO_2 , they are believed to lie along opposite directions (+ and -) in one sheet of atoms and along directions at right angles to these in the next sheet. Or there may be even more complicated arrangements: the important point in anti-ferromagnetics is that there is internal compensation.

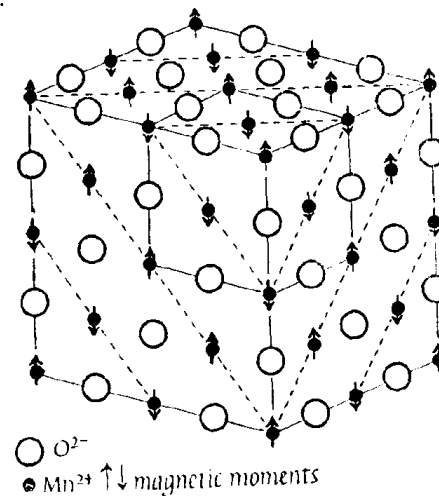


Fig. 3. Chemical and magnetic unit cells of manganous oxide, MnO .

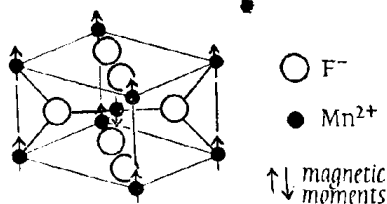


Fig. 4. Chemical and magnetic unit cell of manganous fluoride, MnF_2 .

From the practical point of view it is probable that ferri-magnetics constitute the most important group of crystals to be studied by neutron diffraction methods, because, as previously mentioned, by the selection of suitable combinations of metals it is possible to obtain cheap materials having almost any desired magnetic properties at any given temperature. It may be that the magnetism of the earth itself is related to the magnetic properties of the minerals that compose it.

It has been mentioned earlier that some rare earth elements absorb neutrons rather heavily. They are also, as their name implies, rather rare, or at least difficult to come by. Nevertheless, both the structure and the magnetic properties of rare earth oxides, such as those of those of neodymium, Nd_2O_3 , praseodymium, Pr_2O_3 , and erbium, Er_2O_3 , are of considerable research interest, especially as the two ions, Nd^{3+} and Er^{3+} , have complementary structures from a magnetic point of view. The magnetic properties of both are derived in part from orbital momentum and in part from spin, but whereas these contributions are additive in erbium ions, in those of neodymium they are mutually opposed. Neutron diffraction measurements on both, made at very low temperatures, may be expected to sort out the two effects. There are indeed preliminary indications that this should be possible. X-ray measurements, apart from giving no information about magnetism, could not even be used to find the positions of oxygen in the presence of the much heavier rare-earth ions, but the intensity of neutron diffraction by the oxygen nucleus is greater than that by Pr and not much less than that by

Nd or Er, so that by using neutrons the crystal structures of these oxides have been completely determined.

The other field of structure analysis in which neutron diffraction is now giving most important information is that of crystals containing hydrogen. Ice is the example that comes first to mind. In water there are twice as many hydrogen as oxygen atoms. X-ray analysis has shown that when water freezes the oxygen atoms are arranged in sheets of puckered hexagons, each sheet being the mirror image of those on either side of it. The distance apart of the O atoms is 2.75 Å. The X-ray methods could not at first locate the H atoms, but spectroscopic evidence has indicated clearly that the water molecules remain as such in ice. In each water molecule there are two H at a distance of less than 1 Å from the O and the atomic bonds make an angle of about 105° with each other. Now there are many ways of arranging such molecules to give the symmetry that ice crystals possess. Neutron diffraction has shown that, in fact, ice (Figure 5) is probably a mixture of all such ways; the randomness that this implies explains not only the intensities and positions of the diffraction peaks, but also the strength of the background intensity.

Many, though not all, hydrogen-containing compounds show some randomness of position of the hydrogen atoms. Where there is a hydrogen bond between two electronegative atoms such as two oxygen atoms, the position of the H forming the bond is near to the middle when the bond is a short one, say 2.5 Å, but is nearer to one end or the other of a long (2.75 Å) bond. Neutron diffraction studies sometimes, however, apparently show *half* a hydrogen atom in symmetrical positions on either side of the centre of the bond; and this must mean that the hydrogen atom is statistically to be found sometimes in the one asymmetric position, sometimes in the other, and may even be moving to and fro along the bond. Where water of crystallization occurs in a crystal, the H atoms are, however, definitely attached to the O atom of the water molecule.

A curious state of affairs occurs in a type of crystal such as potassium dihydrogen phosphate, KH_2PO_4 , which shows peculiar electric properties. These crystals are called 'ferro-electric'

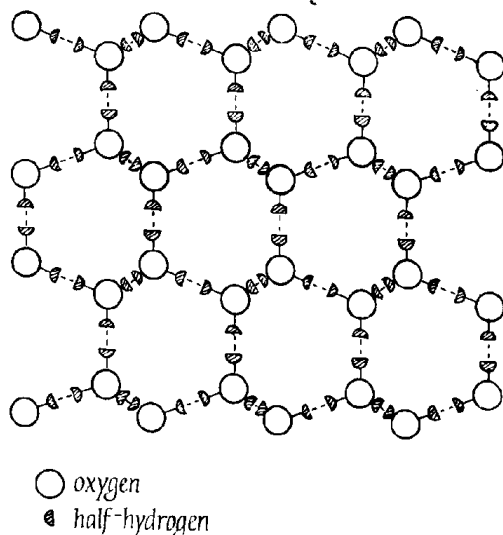


Fig. 5. Apparent structure of ice, seen parallel to the (0001) puckered hexagonal sheets, looking along the a axis.

because their electrical behaviour is analogous in many respects to the magnetic behaviour of such crystals as iron itself. At a certain temperature, or within a given range of temperature, there is some ordering of the crystal structure which causes it to possess directional electrical properties. Compression along a certain crystal direction, for example, causes the appearance of a considerable positive charge at one end and of an equal negative charge at the other end of that particular crystal axis: in other directions this 'anomalous piezoelectric effect' does not occur. Neutron diffraction has shown that outside the temperature region where the large directional piezoelectric effect occurs there is a randomness of position of the H atom along all hydrogen bonds in the KH_2PO_4 crystals; but that if, within the ferro-electric temperature range, KH_2PO_4 crystals are suitably placed in an electric field, the H atoms move definitely to

one end of the hydrogen bonds and stay there, reversing their positions if the field is reversed.

One ingenious device used to determine which neutron diffraction effects are due to H is the replacement of part of the hydrogen by deuterium. Ordinary hydrogen has a nuclear scattering amplitude for neutrons of $-0.38 \cdot 10^{-12}$ centimetre, deuterium has one of $+0.65 \cdot 10^{-12}$ centimetre. A suitable ratio of the two isotopes will result, therefore, in the complete cancelling out of the coherent interference effects due to both. Such a nihilistic substitution method is only possible for neutron diffraction, for which negative as well as positive scattering amplitudes exist. It can also be applied in certain cases of alloys or inorganic compounds where, for instance, $\text{Cr}(+0.35)$ and $\text{Mn}(-0.37)$ could be balanced against each other.

A number of heavy metal hydrides and deuterides, including uranium deuteride, thorium hydride and deuteride, and zirconium deuteride, have been examined by neutron diffraction methods (Figure 6), and the positions of the light atoms determined by what is known as the 'trial and error' method. This means that the diffraction effects to be expected on the basis of likely structures are calculated, and that structure is taken as correct which accounts as closely as possible for the observed data.

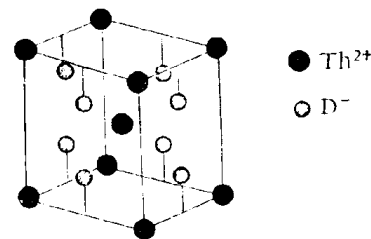


Fig. 6. Unit cell of thorium deuteride ThD_2 .

In the case of somewhat more complicated molecules such as oxalic acid dihydrate, or sodium sesquicarbonate dihydrate, X-ray diffraction has already given a detailed electron density

distribution map of the structure. In some cases the latest and most refined X-ray methods can even indicate H electron contours and probable thermal vibrations of individual atoms. It is then possible, beginning from the atomic positions given by the X-ray work, to use the observed neutron diffraction data to carry out an independent analysis. On the map so obtained the positions of the atomic nuclei will be seen surrounded by diffraction rings due to thermal vibration effects and to the fact that some reflections which cannot be measured, because they would correspond to values of 2θ greater than 180° , are nevertheless not of negligible intensity. In spite of such diffraction rings, neutron maps may be expected to give the positions of the hydrogen nuclei much more precisely than X-ray maps can do. In any case, the two methods, X-ray and neutron diffraction, when used at their highest level of accuracy, may be expected to be complementary to each other, the one giving information about the electrons, the other about the nuclei of the atoms and also about the magnetic moments, where these exist.

FURTHER READING

A good, recent, and much fuller account will be found in:

Bacon, G. E. *Neutron Diffraction*. (Monographs on the Physics and Chemistry of Materials.) Oxford, Clarendon Press, 1955.

NATURE CONSERVATION AND ECOLOGY

J. B. CRAGG

But the resources of the earth in food, soil, water, energy, metals, chemicals, and a dozen other necessities of the new progress, are not unlimited.

A. V. HILL

I LIVE in County Durham. It is a beautiful county, but its beauty is barely glimpsed by those who only see it from a London train bound for Newcastle or Scotland. Instead, they see its pit-heaps and open-cast mountains. They may notice its magnesian-limestone escarpments heavily scarred with quarries. The coast line does not come within their vision. If it did, they would see a beach at Blackhall famous in the nineteenth century for its fauna and now coated with colliery waste. They would notice a 'Lido' spreading its way along the coast soon to envelop an area which is one of the last localities for rare hybrid forms of the Brown Argus butterfly (*Aricia agestis*). With sights of this order one might be justified in borrowing Bill Walker's tongue and saying 'Wot prawce selvytion nah?' But these happenings in Durham County are to some extent taking place everywhere. The problems raised, whilst they concern scientists, are economic and moral and cannot be resolved by scientific action alone. Are we to preserve the fauna and flora of the Durham magnesian limestone at the price of robbing the steel industry of readily obtainable dolomite? As a nation can we conserve our natural resources and yet maintain living standards that are wedded to material production? Obviously, improvements in industrial efficiency, in the utilization of wastes, economies in factory layout, limitations to the spread of towns, are all steps towards

nature conservation. It is important that this is realized, for the work of the biologist in preservation can be rendered useless by the operation of a new factory or the invasion or division of a habitat by a sewage works or a road.

The biologist, concerned with his own special aspects of this broad problem of preserving natural resources or seeing that they are used as efficiently as possible, has one great advantage over the production engineers, civil engineers, chemists, and mining technicians. The material with which he deals has the capacity, if used in the right way, of self-reproduction. Thus, instead of dispersing or destroying his raw materials, he should be able to ensure their self-perpetuation.

Three things should be the concern of the biologist in nature conservation. There is need for the preservation of tracts of land for mental and physical enjoyment, health, and recreation. Some areas should be maintained for their intrinsic scientific value where rare species can be preserved or where natural or semi-natural country can be kept in being. Finally, in the vast field of applied biology, largely because of a rapidly increasing world population, there is greater need than ever for a thorough understanding of the ways in which the soils and seas of the world can be utilized with ever-increasing efficiency.

These three aspects are not mutually exclusive. To maintain a balanced agriculture and silviculture is to maintain much of Britain's beauty. To find out how such areas have remained fertile for hundreds of years and to study in detail the occurrence of rare species or natural assemblages of organisms should give us the kind of scientific information – ecological information – which can form the basis of sound practice in the future.

THE CLIMAX CONCEPT

In a given set of environmental conditions animals and plants live together in a fairly settled way. This apparent condition of equilibrium is for the most part self-sustaining and will go on for many years. Under lowland conditions in many parts of England, oakwoods with their associated animals represent the natural climax or stable condition. On some of our mountains

and hills, moorland is a readily recognizable climax. The preservation of climax types should represent a fairly easy problem for the ecologist provided that he has an adequate amount of territory made available to him and provided that the area is in a fairly natural state when he takes it over. But these conditions rarely hold, and our oakwoods, on which British naval supremacy was supposed to rest, show this only too well. Even in the sixteenth and seventeenth centuries oak was among our major imports, so seriously had our natural oakwoods been depleted. In parts of Britain the day has long passed when the oakwood climax could maintain itself: for the simple ecological reason that interference with one part of a climax can bring about not only a disturbance in the equilibrium position but set up irreversible changes. Thus the indiscriminate killing of owls, hawks, jays, stoats, and weasels in order to protect unnaturally high populations of game animals has resulted in increased numbers of mice, voles, and rabbits. These rodents, by nibbling off the tops of seedling trees, have prevented natural regeneration of woodland. Even where young trees have been planted their roots have been severely damaged and growth impeded. The surrounding ground vegetation is similarly affected, but in most circumstances vole activity seems to check the growth of some plants, allowing other species to come in, thus producing a more varied ground cover.

The results that can follow from interference with a climax formation have been graphically described by A. S. Leopold and F. Fraser Darling. They were asked to carry out a survey of the large game animals in Alaska in order to outline the present status of the species and to examine the problems of their conservation.

The three major species on which the local populations depend are caribou, reindeer, and moose. The caribou from being a common animal throughout Alaska is now considerably restricted in its distribution (Figure 1). In the early 1920s the population was estimated at approximately one to two millions; now it is somewhere around 130,000. The investigators have concluded that the major factor responsible for this decline has

been fire damage to the lichens which constitute the main winter-feed of these animals. The lichens reach their maximum abundance in areas of sparse spruce forest which form the vegetation climax of certain parts of the Alaskan tundra. Under the extreme climatic conditions which exist there, burning and clearing have created havoc with the vegetation. The destruction may have gone too far in some areas, and even in favoured zones it is estimated that it may take 50 to 100 years for the climax to recover. Where felling or burning have removed the spruce, a sub-climax cover of willow, aspen, and birch scrub has developed. Under normal circumstances this would give cover to the slowly regenerating spruce. This scrub forms the main food of the moose. Thus with decline in caribou populations there has been a consequent rise in numbers of moose.

Burning and clearing have not been the only agencies at work. Whalers, walrus-hunters, and mining development in the nineteenth century so depleted the natural food resources of the local Eskimos that many died from starvation. Reindeer were imported to make up the lack of available food in the coastal regions. Approximately 1,300 animals were introduced between 1891 and 1902. By 1932 their progeny numbered over 600,000. The reindeer uses the same vegetation as the caribou, but being a more tractable animal it was herded by man. The local Eskimos, with a very different way of life from that of the nomadic Lapp, and without their long experience of reindeer herding, allowed the introduced reindeer to overgraze its range to such an extent that numbers crashed in the 1940s, falling to 27,000 in 1952 (Figure 2).

Even in circumstances where man does not interfere with a climax, the equilibrium condition is certainly not fixed. There are years when some organisms in the climax are abundant and at other times the same species may almost disappear. The greatest biological stability appears to result when a large number of plants and animals are involved in making up the climax. Thus natural tropical forest with its vast numbers of species will probably show greater stability than a climax developed in sub-arctic or arctic regions. Nature is never static and the durability

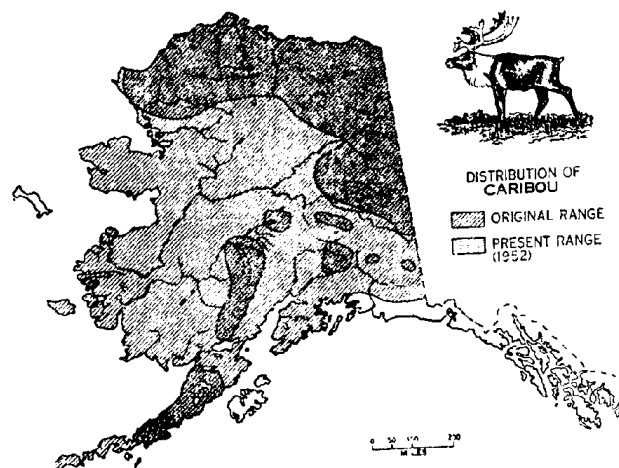


Fig. 1. Original and 1952 ranges of caribou in Alaska, showing shrinkage of area occupied in central and southern regions.

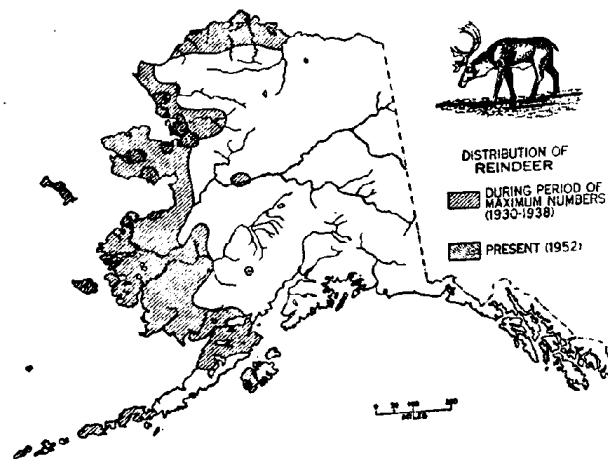


Fig. 2. Area occupied by reindeer at period of maximum numbers (1930-38) and 1952. (Both figures from Leopold and Darling.)

of the climax depends largely on the time scale which we use when viewing it. The past history of our northern moors shows periods when forest cover reached to the 2,000-ft. contour. Long-term fluctuations in climate alter the constitution of a climax. Thus, there are distinct signs of major changes taking place in the climate of north Europe at the present time. Already, a northward spread of many species of birds and mammals has been recorded in Finland. These climatic changes complicate the story of the caribou outlined above and they make the task of the ecologist concerned with nature conservation a difficult one.

CYCLES IN ANIMAL POPULATIONS

Although the climax represents a condition of relative stability, yet within a climax, or in circumstances where no obvious climax exists, many animal species show marked fluctuations in abundance. The animal ecologist has given a considerable amount of attention to the possible causes of such fluctuations. When he is in a position to forecast the population trends of an organism he will then possess information on which sound conservation measures can be based. One aspect of this work which has a particular bearing on problems of nature conservation concerns studies on alleged cyclic fluctuations in the numbers of some animals living under natural conditions.

Most of the species which have been studied are rodents, fur-bearing carnivores, predaceous birds, and game birds. In the past it has been too readily assumed that their fluctuations in numbers follow a very regular cyclic pattern, and that fluctuations in one species may be closely linked with those of another species. The latest comprehensive study on these phenomena is a series of papers published in the *Journal of Wildlife Management*, 1954.

The major cycles which appear to exist are of 3 to 4 years' and 8 to 10 years' duration. The short type occurs in lemmings, small grass-eating rodents common on the arctic tundra, the varying hare or snowshoe rabbit, musk-rat, and certain birds such as the ruffed grouse which is a common game bird in North

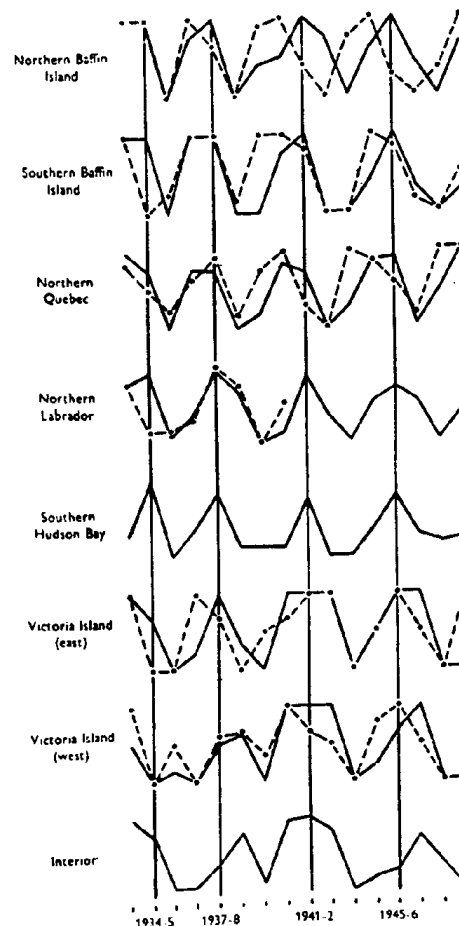


Fig. 3. Reports for 1933-49 of snowy owl migrations and of changes in numbers of lemmings and arctic foxes in northern Canada. The percentage of observers reporting increase or abundance is shown by broken lines and dots for lemmings, and by solid lines for foxes. Vertical lines are years of major migrations of snowy owls. (Helen Chitty.)

America. The 8-10 year cycle is characteristic of major predators such as the lynx and, according to some authorities, the arctic fox and mink. In some cases it is possible to recognize that fluctuations in the abundance of a rodent coincide with similar fluctuations in animals preying upon it. Thus in northern Canada the synchronization in numbers of lemming and arctic fox is brought out in Figure 3. The vertical lines on the figure represent the major southerly migrations of the snowy owl, and it will be noticed that these occur when lemming populations are on the decline. The owl is affected in another way by the disappearance of its food supply. When food is plentiful 6 to 8 chicks per nest reach maturity, but when greater areas have to be covered in the search for food only 2 to 3 survive, and when food is very scarce the owls may not attempt to nest. Thus available food material affects not only the predator population existing at the time of the 'crash' but also the succeeding generation.

A considerable amount of attention has been given to the possible causes of these fluctuations. On the assumption that there was a marked synchronization of the fluctuations occurring in different species, it has been variously suggested that they are related to the sun-spot cycle, to the amount of ozone in the atmosphere, which in turn is correlated with the amount of ultraviolet radiation reaching the earth, and to the amount of ultraviolet radiation as such. These suggestions have been made in relation to very different organisms. Detailed work has shown, however, that the individual cycles, that of the lynx in particular, do not fit in with sun-spot activity; and, apart from our inadequate knowledge of variations in ozone content and the amount of ultraviolet radiation at sites where the animals live, there is sufficient evidence in the case of some birds and mammals to indicate that fluctuations in numbers, even in one species, are not in phase over the greater part of its range. It is unlikely therefore that an extra-terrestrial agency is the master factor causing these fluctuations.

In 1949, P. Palmgren, a Finnish zoologist, came to the conclusion that the 3-4 year fluctuations resulted from random variations, perhaps in climatic factors, operating on population

densities which are in themselves dependent not only on the climate of the year in which they occur but on the numbers carried over from previous years. LaMont Cole has carried Palmgren's analysis a stage further and has concluded that much of the periodicity data now available is consistent with the hypothesis that population changes reflect purely random fluctuations in environmental conditions. He is very critical of the way in which minor peaks have been neglected or data smoothed. Using a set of random numbers as his basic material, he has shown that by taking averages a random distribution of the numbers takes on the appearance of a set of population fluctuations (Figure 4). J. J. Hickey, an ecologist of the Department of Wildlife Management, University of Wisconsin, has studied much of the data available for Canada using Cole's mathematical approach. While agreeing with Palmgren and Cole about the random nature of short-term fluctuations, he concludes that the 8-10 year type of periodicity is not a product of random variations in the environment. Of the cases studied by him, only the lynx emerges as possessing a definite non-random type of cycle, but it may hold for two other predators, the red fox and mink. In the case of non-predators, Hickey's own analysis of the Hudson Bay Company's returns for the 1847-1903 and 1931-49 periods indicate that the snowshoe rabbit had a definite short-term periodicity with averages of 4.6 to 3.8 years respectively. MacLulich concluded that a long-term periodicity with a mean of 9.8 years existed in this species. Hickey does not think on present evidence that this claim for a non-random type of periodicity in this species should be disregarded.

There is always a danger that the biologist faced with a mathematical analysis of animal numbers will concentrate too much on the numbers and not give enough attention to the organisms which they represent. It was Hickey who said of research on cycles '... that ecologists rush in where mathematicians fear to tread'. Much of the confusion that surrounds discussions on cycles and oscillations arises from the very doubtful nature of many of the records. Much of the information utilized for cal-

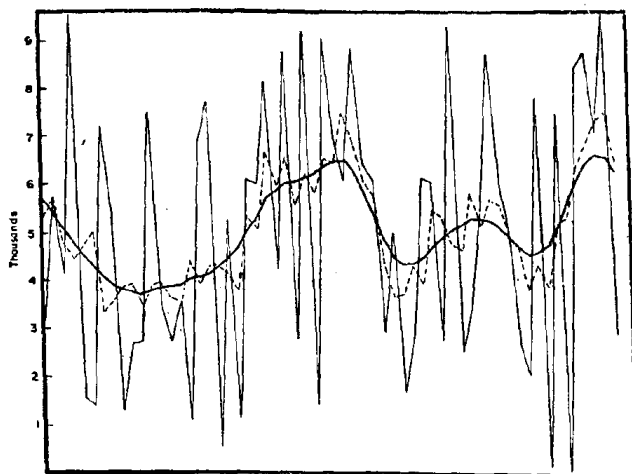


Fig. 4. Cycles produced by smoothing random data. Tippet's random numbers (fine line) plotted in order of appearance, and derived from these a simple 5-point moving average (broken line) and a triple 5-point moving average (heavy line) of the same numbers. (From Lamont Cole.)

culating fluctuations in mammals has come from the Hudson Bay Company's records. These consist of pelts handled by them in the course of trading. Such records must be influenced by a wide variety of environmental and social factors.

The great need in this branch of ecology is detailed knowledge of the individual species, their habitat requirements and their behaviour patterns. The value of such studies becomes obvious when considering the work of P. L. Errington and D. H. Chitty.

Errington, who has worked on musk-rat populations in Iowa for some 20 years, has demonstrated that, when the snowshoe rabbit and ruffed grouse populations in northern Minnesota and Wisconsin were on the increase, his musk rats, miles away in Iowa, reached high densities without the occurrence of intra-specific strife. On the other hand, when the rabbits and ruffed

grouse numbers were falling, the musk-rat densities also fell. What is perhaps more important, although densities were lower with no apparent change in the amount of food available, intra-specific fighting among adults was common, and even helpless, young, individuals were sometimes attacked. Changes in the musk-rats' resistance to disease were also noted. During the rise of snowshoe rabbit and ruffed grouse populations, musk-rats were highly resistant to a haemorrhagic disease. Coinciding with the collapse of the rabbit and grouse populations, the musk-rats' resistance to disease declined, and large numbers were killed off. As Errington points out, the musk-rat is semi-aquatic and differs markedly in habitat requirements from the two other animals, yet there was an apparent synchronization in their rises and falls.

It is the variations in behaviour and resistance to disease which are probably the most reliable and at the same time most striking feature of Errington's studies. They link up closely with the results obtained by Chitty when studying the common short-tailed vole in the Lake Vyrnwy region of Wales. He concluded that intra-specific strife was a major factor in bringing about the collapse of the population but that its action was indirect. His interpretation rests on two assumptions: that physiological changes occur among crowded animals, and that their descendants are adversely affected by the condition of their parents. The phenomena recorded by Chitty and Errington have a great deal in common, and we must look forward with interest to the final conclusions of Chitty and his Oxford colleagues.

Whether the changes in behaviour are themselves 'triggered off' by environmental factors is not known. D. Lack, who has worked for a long time on bird populations, is of the opinion that in many species the fluctuations are the direct result of food shortages. In some cases, a species may be affected by the decline in another species' food supply. Thus, Lack explains the fluctuations in certain ruffed grouse populations by suggesting that predators such as Red Fox, Horned Owl and Goshawk make more inroads on ruffed grouse numbers when their normal food, the snowshoe rabbit, is on the decline.

To summarize modern work on cyclic phenomena in the numbers of animals, it is evident that many of the fluctuations are random in origin and possible synchronizations between species are not necessarily caused by the operation of common factors. As for the remainder, the fluctuations of a non-random nature, the vast majority of workers do not attribute their occurrence to the actions of extra-terrestrial factors. The possibility of genetic changes influencing population densities has barely been considered, and it is on detailed experimental studies that future progress would seem to depend.

THE CONTROL OF ANIMAL POPULATIONS

Cyclic fluctuations are only a special aspect of the population picture. In temperate and tropical countries oscillations of the type just described are not observed, and the falls and rises appear to be wholly random, at least for the populations where data are available. There are few organisms for which we have quantitative records stretching over a long period. Schwerdtfeger's data on certain pine forest pests covering some 80 years are exceptional both in length and quality. In the majority of cases we must rely on the qualitative observations of competent naturalists. It is generally agreed that animal and plant populations do not go on increasing without pause. Instead, the size to which they can grow is usually restricted. Some ecologists believe that limitation is brought about by factors which are density-dependent; that is, by factors which operate more intensely on average members of the population the greater its density. This aspect of population study has been dealt with by J. B. S. Haldane in *New Biology* 15 and by F. S. Bodenheimer in *Science News* 30.

The most clear-cut cases of density dependence are those where a parasite is wholly confined to a single host species and where the life cycle of the parasite is closely adapted to the life cycle of the host. Thus the joint studies of De Back and Smith on the Hymenopterous parasite (*Mormoniella vitripennis*) which attacks the pupal stage of the house-fly (*Musca domestica*) demonstrates the way in which the populations of host and para-

site are affected (Figure 5). In addition to showing that oscillations occur in the populations, they demonstrated that there was close agreement between their results and the results to be expected on the basis of the mathematical treatment by A. J. Nicholson and V. A. Bailey. But there is a great jump from the laboratory experiment carefully planned to fit the mathematician's assumptions and the state of affairs which prevail under field conditions.

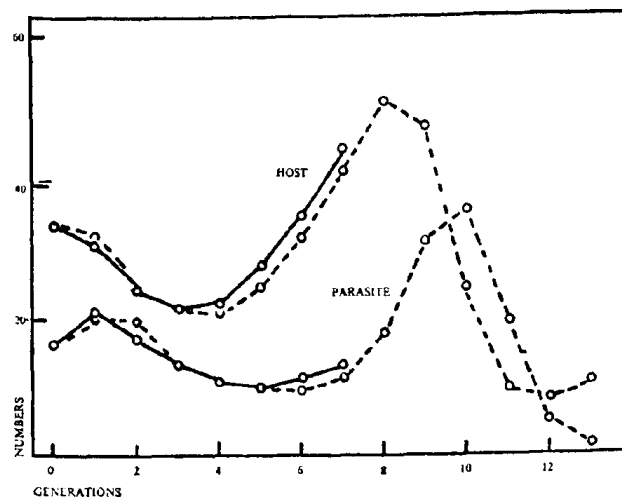


Fig. 5. The fluctuations in numbers of a parasite (*Mormoniella vitripennis*) and its host (*Musca domestica*). Solid lines show the experimental results and broken lines the results expected on the basis of Nicholson and Bailey's mathematical model. (After de Bach and Smith.)

In the future the conservation of an animal species will no doubt be based on principles derived from a knowledge of the factors which regulate the numbers of animals. At the present time, the information must be sought for each individual species and there are very few cases where it has been collected on a

truly quantitative basis. In fact *The Natural Regulation of Animal Numbers* by Lack, published in 1954, was the first attempt to cover this subject in detail. There is great scope for research in this field and, in particular, a great need for routine investigations on a mass-production scale.

Too many of the hypotheses concerned with the regulation of a species have been derived from over-simplified mathematical models in which insufficient allowance can be made for the ever-changing conditions which beset an organism in nature. Information is therefore wanted on the effects of a fluctuating environment on an organism throughout its life-history under conditions where the fluctuations can be accurately measured and where the interactions of different species can be studied in detail. Experimental work of this type will require large numbers of technical assistants and the development of much self-recording apparatus. The approach is of course being developed to some extent in applied biology. Various insecticide manufacturers have developed the routine testing of insecticides on a big scale, but the number of test species and the variety of conditions studied are relatively small. In the study of 'stored products' insects, particularly those associated with grain, great strides have been made in the routine analysis of the effects of changes in the environment.

CONSERVATION PROJECTS

The call for conservation most often arises when some region or organism has been so exploited by man that the region becomes biologically useless or the organism approaches extinction. Generally speaking, such conditions arise because the natural climax condition has been broken down and man has failed to simulate its essential features when cropping the land or removing the animals be they small game or bison.

Erosion or soil loss (Inset 8) is a natural process brought about by the action of climatic factors, such as wind and rain, on the land surface. Under stable conditions the natural vegetation protects the land against erosion and new soil is formed in sufficient quantity to make good any loss. When the vegetation is

destroyed, soil loss may become excessive. Thus, on a forest-cleared slope, heavy rains will scour deep channels and wash away top soil which has taken years to accumulate under the protection of the trees. A large number of factors influence the extent of soil erosion: rain, wind, soil-type, topography, vegetation-cover, amount of grazing, type of stock, methods of agriculture, the economics of land-values, all may be involved in the erosion problem. The dust clouds containing the top soil of far-western farms which in 1935 swept over huge areas of the United States, fortunately reaching the centres of government and power in Washington and New York, were a dramatic demonstration of the dangers which can follow bad land use. Israel, faced with tremendous problems of land reclamation (Inset 6-7), is approaching them on a broad basis. So advanced is their approach that they actually have a Government Chief Ecologist, Dr H. Boyko, who has recently described the problems facing plant ecologists concerned with increasing productivity under arid conditions.

As Boyko points out, something like a third of the earth's surface lies within the arid zone and much of it has become semi-desert through man's activities. The first task of the Israel scientists has been to determine the natural climax associations of the areas to be reclaimed. In many cases, the best fodder plants have disappeared through the effects of thousands of years of over-grazing. Search has therefore been made for rare species in places practically inaccessible to the grazing animal. The fencing off of areas is also a method which may allow some natural regeneration to take place. Thus, in southern Israel a small shrub, *Artemisia herba alba*, is avoided by animals and as a result has colonized vast regions. When areas are protected against grazing animals, previously rare grasses have appeared in profusion and in a matter of several years have replaced the *Artemisia*.

Over-grazing is not only a problem of the land. In various parts of the world man has functioned directly as a predator or a grazing animal on both fish and marine mammals. He has jeopardized the economic life of many of his fisheries and has

rendered extinct or almost extinct interesting and even commercially valuable marine mammals.

The damage that has been done to fisheries throughout the world has been well described by E. S. Russell in *The Overfishing Problem*. The solution, a difficult one to operate, lies in regulating the fisheries so that adequate mature stocks are maintained to preserve the fishing. Although the total weight of fish taken may not fall when the intensity of fishing increases, a change takes place in the age and size composition of the catch, younger and smaller fish predominating. In fact, the young forms may show more rapid growth because of the reduced density of the stock. Thus young plaice 'transplanted' from the practically unfished and overcrowded waters of the Barents Sea to the North Sea grew at a faster rate than they would have done under Barents Sea conditions.

Seals, both in the northern and southern hemispheres, require protection. Soon after its discovery in 1775 by Captain Cook, South Georgia became a centre of intense sealing activities. As early as 1822 the southern fur seal (*Arctocephalus australis*) was probably extinct. Towards the end of the nineteenth century the Elephant Seal (*Mirounga leonina*) had been reduced to such an extent that the sealing industry practically died out. A threat of commercial loss is often a means of introducing conservation measures, and in 1899 restrictions were introduced. At present the killing of certain species is forbidden, sealers must hold a licence, and they are allowed to remove only a certain number of seals per year. Dr R. M. Laws, a biologist on the staff of the Falkland Islands Dependencies Survey, has recently surveyed the South Georgia seal population. In 1951 he estimated the Elephant Seal numbers of the island to be about a quarter of a million. This might appear to be a healthy state of affairs, but on closer examination the situation is far from satisfactory. On the South Orkney Islands, bulls have an average age of 20 years, whereas on South Georgia Laws reported that 'not one of 100 bulls examined reached an age of more than 10 years . . .'. The age of those killed by sealers ranged from 2 to 10 years. He found evidence that the stock of bulls is diminish-

ing, and the fall in oil output is a fairly sure sign that smaller animals are making up the 'kills'. Because there are fewer males the harem size has been increased, and this in time will probably mean that not all cows will be fertilized. It is of interest to note that the species is reacting to the activities of man the predator, for there is some evidence that the Elephant Seals are leaving the larger bays for more inaccessible ones. New conservation measures based on these recent biological studies are being introduced. Only males above a certain size are to be killed and some 10 per cent of the bulls above this size are to be left in the herds. These scientifically based restrictions are only in their early stages, but there is every likelihood that they have been introduced in time. If only some similar organization had been prepared to take conservation measures in the Arctic before even South Georgia had been discovered, Steller's Sea Cow might have been saved not only for its scientific interest but perhaps for human food. Over two tons in weight, grazing, without fear of man, on seaweed in the region of Bering Island, the species was completely destroyed within 27 years of its discovery.

Conservation now seems to be the rule in the extreme regions of the world. The Danes have a number of laws aimed at conserving the musk-ox in north-east Greenland. Trappers may in each season kill up to six, of which only two may be cows, and persons other than trappers are only allowed to kill one per season. Ships' crews, which in the past have been responsible for the total destruction of many species, are limited to two musk-oxen per vessel and they may only be hunted by members of the crews. One clause in the Danish regulations applying to N.E. Greenland might be given some consideration in other countries; gin-traps and poison are forbidden.

Limiting the amount of killing is not the only method available for the conservation of marine or game animals. In the case of caribou, Leopold and Darling stressed the importance of conserving the winter range. This aspect of conservation cannot be overstressed for it is on matters like this that the ecologist, because of his outlook, should prove to be of particular value.

The work of Leopold and Darling should be read by all interested in conservation problems. It shows what two people with an 'eye for country' and possessed of wisdom born of long experience in the field, can achieve in the small space of four months. Obviously there are many problems concerning caribou reindeer and moose which remain to be studied in detail. But their work is an excellent example of an ecological reconnaissance carried out to pave the way for scientifically based conservation measures.

CONCLUSION

The ecologist's role in nature conservation is that of a technician. Fundamentally conservation is not a scientific problem but a moral one. It depends on the wide acceptance of the principle that we must make ourselves responsible for the proper use and maintenance of the resources of the world.

It would be wrong to imagine that without ecological information there would be no nature conservation. The great bulk of nature conservation in the world today owes little to the activities of the ecologist. A considerable fund of knowledge has been built up by farmers, foresters, huntsmen, estate managers, fishermen, and naturalists. It is the ecologist's job to study the results of their work and so determine the reasons for the survival of the communities for which they have been responsible.

It is not enough for the ecologist to look at the present state of a plant or animal association. In order to plan or forecast its future development he should know something about its biological past. One of the functions of bodies such as the Nature Conservancy must be to initiate experiments planned to continue for hundreds of years, so that long-term changes in plant and animal populations can be followed with accuracy. There have been many lost opportunities for the collection of such information in the past. Charles Darwin discerned this need when in 1863 he wrote to Dr Julian von Haast, a German by birth, settled in New Zealand: 'I would venture to suggest to you to urge on some of the capable members of your institution to observe annually the rate and manner of spreading of European

weeds and insects, and especially to observe *what native plants must fail...*

A region of the world that cries out for the attentions of the ecologist in the role of nature conservator is the Tropical Rain Forest belt. In the 1930s it was estimated that evergreen tropical forest, the characteristic climax of the wet tropics, constituted something like half of the world's forest area. The invasion of the tropics by Western man has so affected the forest that P. W. Richards states: 'Unless determined efforts are made to halt the destruction, the whole of the Tropical Rain forest may disappear within the lifetime of those now living.' The forest has existed for many millions of years. With its vast wealth of species it has acted as a reservoir of genetic diversity and much of the flora of the temperate regions has been derived from it. As Richards points out, the damage wrought by man during the last one hundred years has probably brought about fundamental changes in the course of plant evolution. It will be a scientific and aesthetic tragedy of the greatest magnitude if nothing is done to preserve large tracts of this primitive forest. The economic potentialities of many of its species are unknown and quantitative scientific studies on its plants and animals barely begun.

I started this article with a reference to County Durham. I did it to impress upon the reader that one of the greatest needs in nature conservation is to control mechanical exploitation. The biologist fares badly as a producer in the midst of an industrial civilization. As Professor W. H. Pearsall, in his presidential address to the section of botany, pointed out at the 1954 meeting of the British Association, the Esso installation at Fawley produces some 50 tons of petrol per acre per day, a modern steel-works can yield 15 tons of steel per acre per day. Against this the best wheat lands can probably be expected to produce just over one ton per acre per year. But you cannot eat steel, and industrial production is in the long-term sense a mechanism for the rapid destruction of natural resources. The conversion of steel into motor cars is not merely a redistribution of steel. In America, according to Fairfield Osborn, the annual car output

also requires the produce of 3,600,000 acres of agricultural land. No doubt some of this loss can be prevented by the use of synthetics, but it is a loss of wool, leather, oil and cotton, to name only the major items, and the fertilizers and manures which have gone to make them, that we are apt to forget.

Louis Bromfield has stated that 'We have had too much thinking and action in terms of money and too little in terms of real wealth'. If ecologists are to make a full contribution not only to conservation for beauty but to conservation for production, then Governments will have to be prepared to spend on ecological and biological research in general far more in proportion to what they spend on other branches of science today.

FURTHER READING

Report of the Nature Conservancy for the year ended 30th September 1954. London, H.M.S.O., 1955.

Biology of Deserts. A symposium edited by J. L. Cloudsley-Thompson. London, Institute of Biology, 1954.

A Symposium on Cycles in Animal Populations. Edited by O. H. Hewitt. See especially for papers by L. C. Cole, P. L. Errington, D. Lack, J. J. Hickey, W. Rowan. *Journal of Wildlife Management*, **18**, 1-112 (1954).

Lack, D. *The Natural Regulation of Animal Numbers.* Oxford, Clarendon Press, 1954.

Leopold, A. S. and Darling, F. F. *Wildlife in Alaska.* The Ronald Press, New York, 1953.

Richards, P. W. *The Tropical Rain Forest.* Cambridge University Press, 1952.

Sears, P. B. *Deserts on the March.* London, Routledge, Kegan & Paul Ltd, 1949.

LIGHTING AND WORK

J. B. COLLINS

LIGHTING is very often taken for granted in these days of complex living, and lighting research is usually thought of only in terms of improvements in efficiency or power of sources. New sources, however, bring their own problems of correct use to add to the fundamental problems of the functioning of this complicated chain of stimulation, registration, transmission, perception, and interpretation, which we call 'seeing'. These problems of lighting and its effect on our efficiency, comfort, and mental and emotional state are being tackled by workers in many different fields, from physicists and engineers to psychologists, and from physiologists to architects. Lighting is a science, for all that art is necessary for the satisfying realization of technical solutions, and, indeed, a sound technical basis is a necessity for the success of any art. The story of present-day research on principles of lighting is, in the main, that of the techniques now being adopted to use the human being as his own meter, to discover what range of conditions best suits him, and of the influence of the results obtained from their use on modern lighting practice.

HOW MUCH LIGHT CAN WE AFFORD?

The most obvious elementary requirement of a lighting technique is to be able to prescribe the quantity of light required for a given visual task. Until recently, the only techniques for determining psycho-physical quantities such as this were those of threshold measurement. The minimum amount of light for an object to be just visible, or letters to be just distinguishable, could be fairly readily determined, but the amount by which the threshold was required to be exceeded for efficient seeing was always a matter of guesswork until some measurement of performance could be devised.

The method of measuring performance which has been used in Great Britain and abroad to obtain the basic data on effect of illumination, was devised by A. W. Beutell, and applied on a practical scale by H. C. Weston. This method combines assessment of speed with that of accuracy, and involves the minimum use of intelligence and manual skill. To assess the visual performance of a subject, he is presented with a card of the type shown in Figure 1. This card contains broken circles printed in black or grey on a white background, with the gap orientated in a random fashion in one of eight directions. The subject is asked to mark (by putting a pencil line through) all the circles which have their gap in one particular direction. He is told to do this as accurately and quickly as possible, and the performance is measured by the time taken to complete the card, and by the number of errors made. The cards are printed in a range of circle sizes, and with inks having a range of densities from black

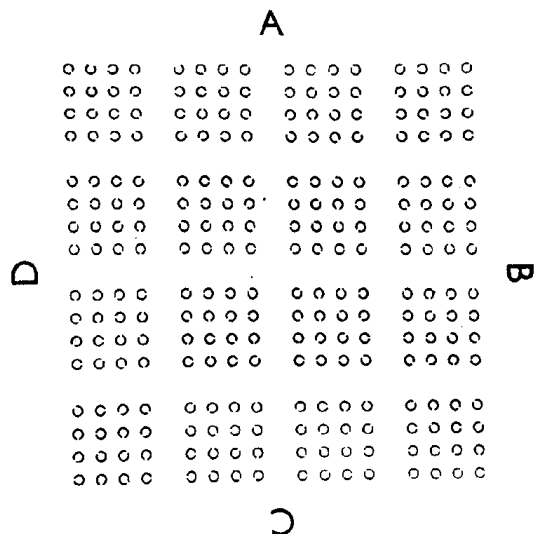
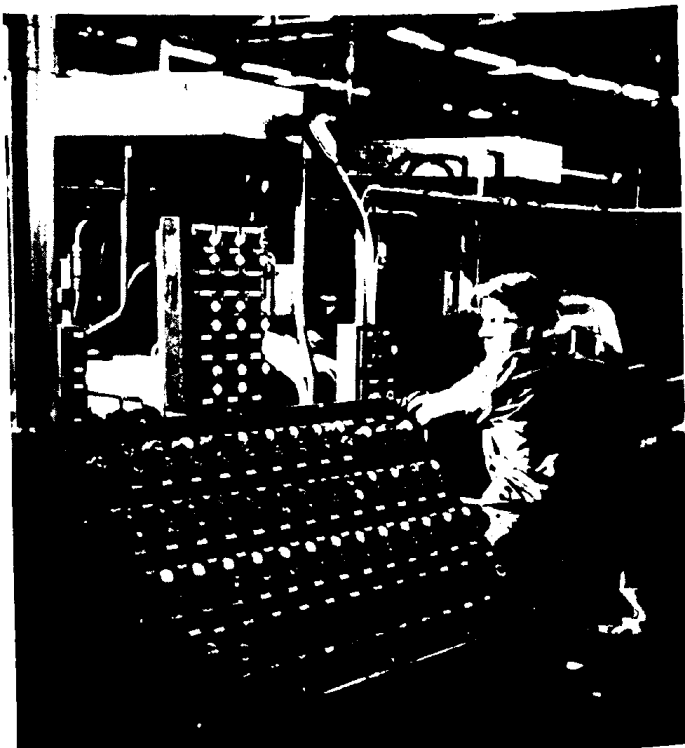


Fig. 1. Weston visual performance chart.

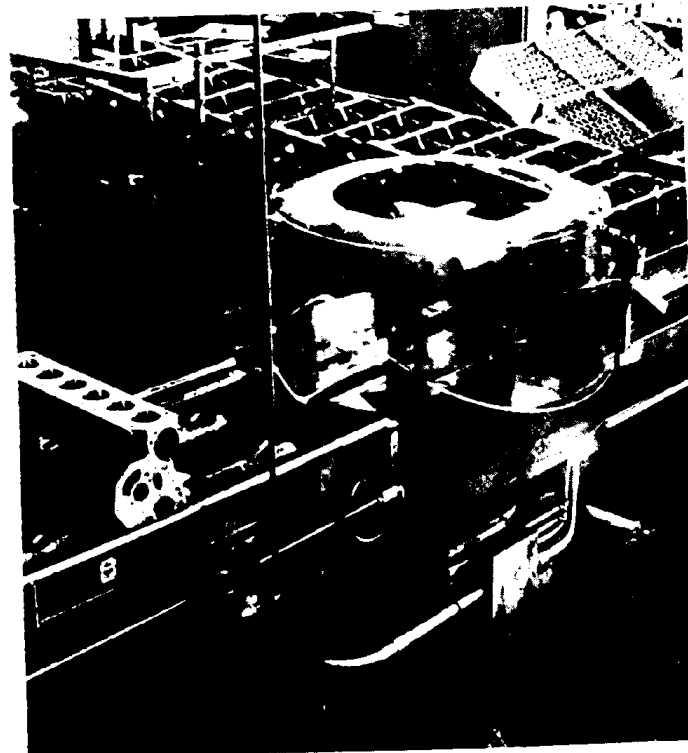
THE AUTOMATIC FACTORY



1. A worker at the Cleveland engine plant changes drills in a Snyder transfer machine that drills 20 holes in a crankshaft from four different directions. Automation moves the crankshafts into place before each battery of drills, starts and stops the drilling action, and sends the crankshafts to their next operation.



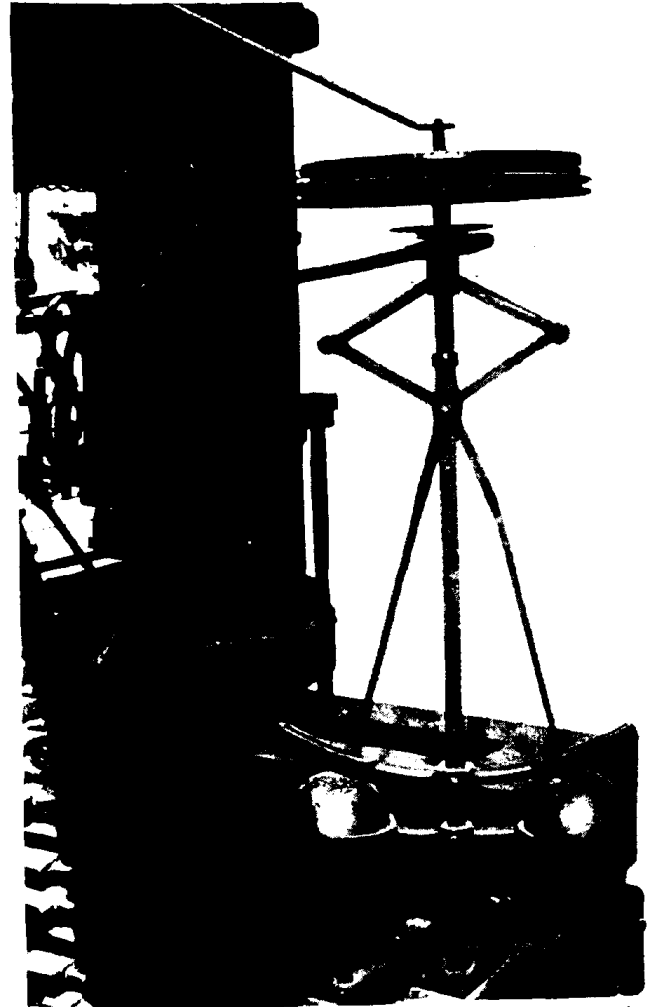
2. This automation control board guides the feeding and unloading of cylinder blocks in a Barnes drill twice as fast as could be done with semi-automatic manual controls. The man takes only an occasional look at the board. Hand labour is much reduced, highly skilled maintenance is increased.



3. The camera catches this automation turntable in the act of whirling a block upside down and turning it through 90°. This changes it from the position on the left to that on the right. (This and the preceding two pictures are by courtesy of the Ford Motor Company, U.S.A.)

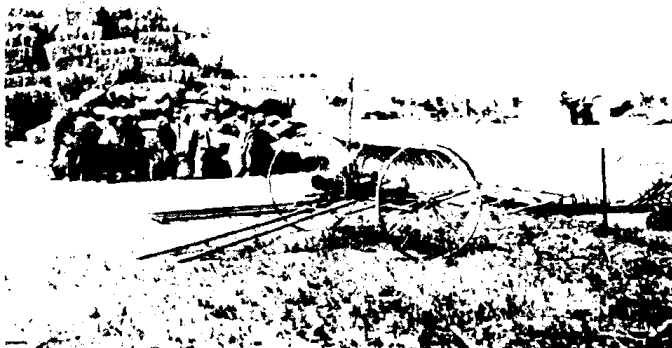


4. Transfer machine showing various stations and the conveyor system. (By courtesy of James Archdale and Company, Ltd.)



5. Speed governor on Boulton and Watt's rotative beam ('Lap') engine, 1788. Variation in the speed of the output shaft alters the position of the lead spheres; through a system of hinged rods this causes a movement of the flanged collar and thus varies the input of steam to the cylinder. (Crown copyright. From an exhibit in the Science Museum, South Kensington.)

NATURE CONSERVATION AND ECOLOGY



6, 7. Land development in the Negev, southern Israel: *above*, a newly established settlement; *below*, after a harvest made possible by water brought from the north. (By courtesy of the Jewish National Fund.)

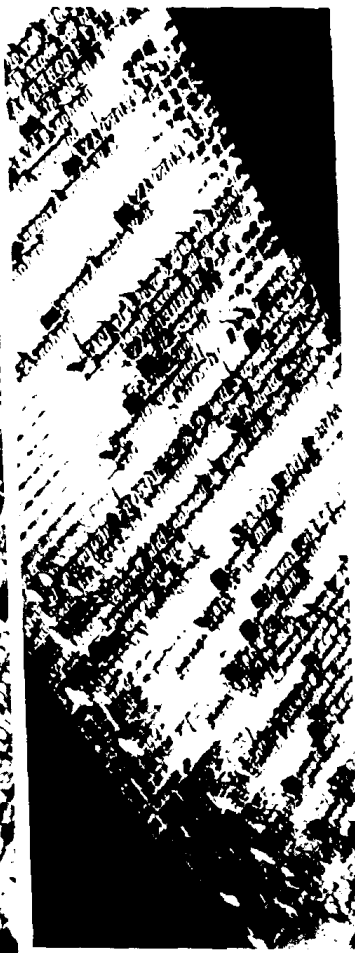


8, 9. Soil erosion and conservation in the Belgian Congo: *above*, erosion in the Ruindi Valley, Parc National Albert; *below*, terrace cultivation at Kibirizi, Astrida-Ruanda Territory. (By courtesy Centre d'Information du Congo-Belge et du Ruanda-Urundi.)

LIGHTING AND WORK



10 (*left*). Set type illuminated by shadowless lighting; note the confusion caused by lack of contrast.

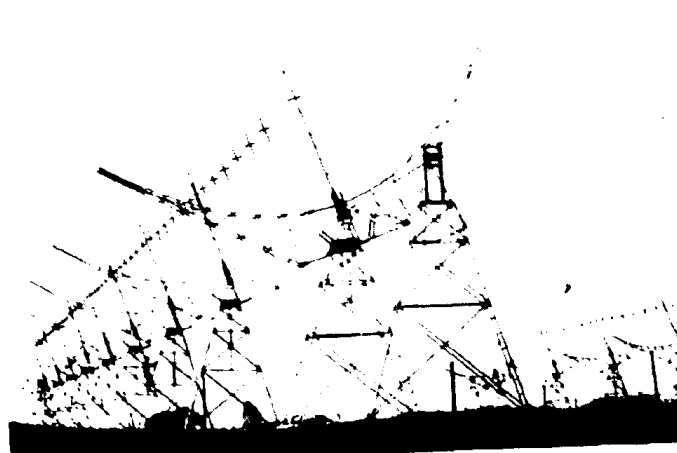


11. Set type illuminated by directional lighting; note improved visibility caused by shadows and specular reflections. (Photographs by courtesy of the General Electric Company, Ltd.)



12. Arrangement of field of view for subjective glare studies. (*Building Research Station, Cambridge, 1931-32*)

RESEARCH REPORT



13. One of four units of the Cambridge radio telescope with which nearly 2,000 'radio stars' have been found

14.



16.



15.



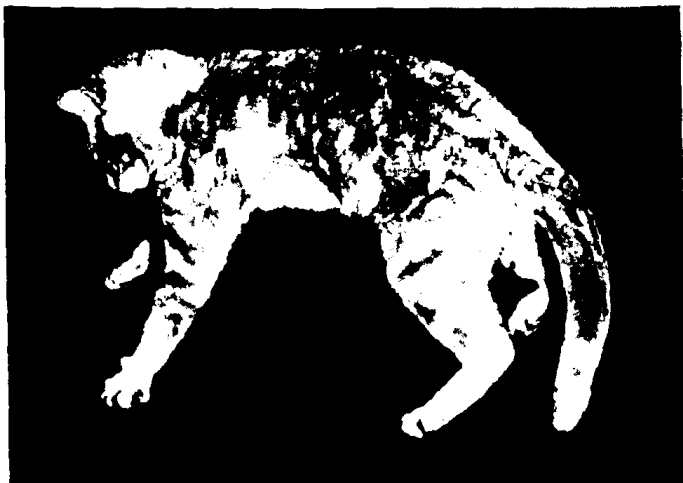
17.



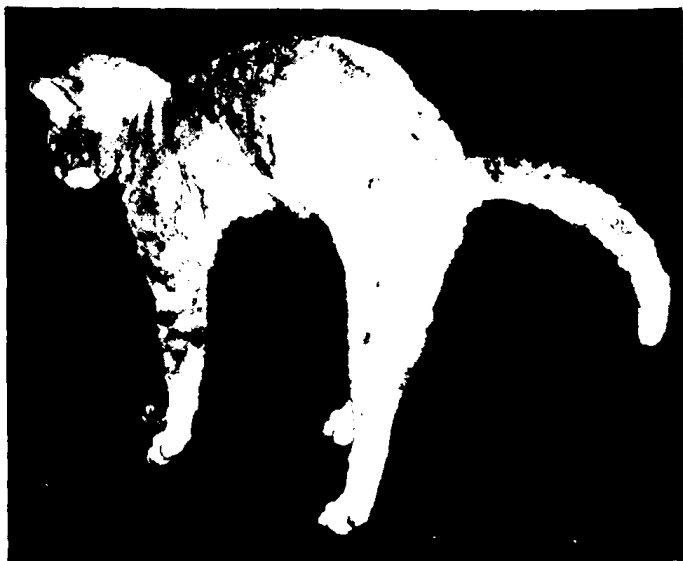
14-17. The photographs on this and the next three pages illustrate the results obtained with an ultra-high speed film taken of the dropping of a cat, held upside down and released. The intervals between the pictures reproduced are about $1/20$ second. The cat's shoulders

turn immediately after release, the tail being rotated in the opposite direction. After dropping only 12 inches, the cat's shoulders have turned through 180° , and, with the tail still contra-rotating, the hindquarters are turned too.

18.



19.



18-21. Continuation of the sequence begun on the previous two pages. With legs extended the cat touches down, then bends its legs to break the fall, the tail being used to maintain balance. Further pictures would show the cat shaking its head, apparently an invari-

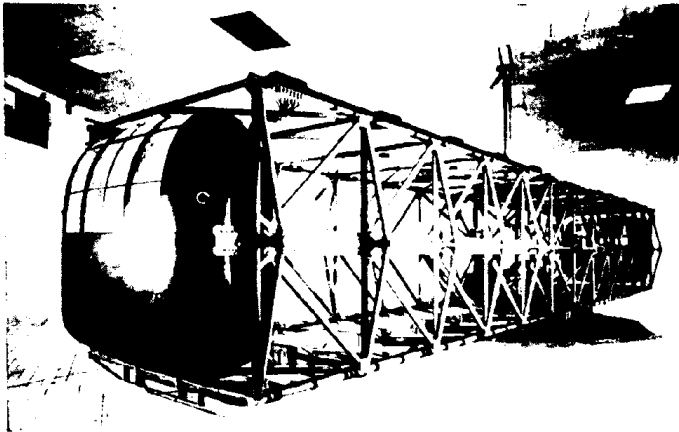
20.



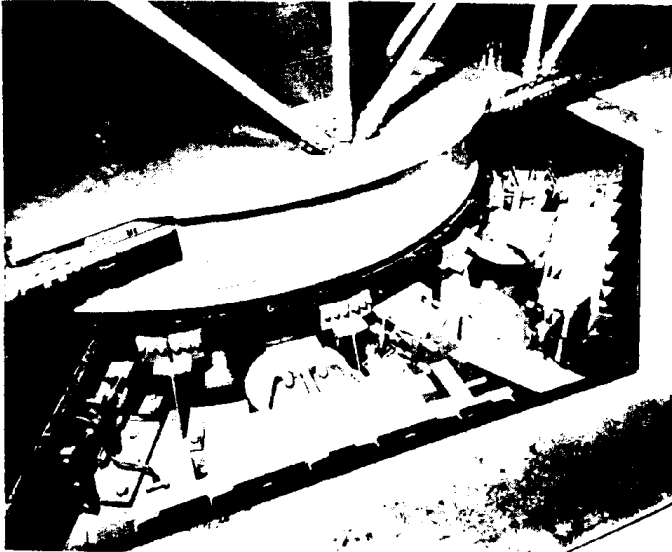
21.



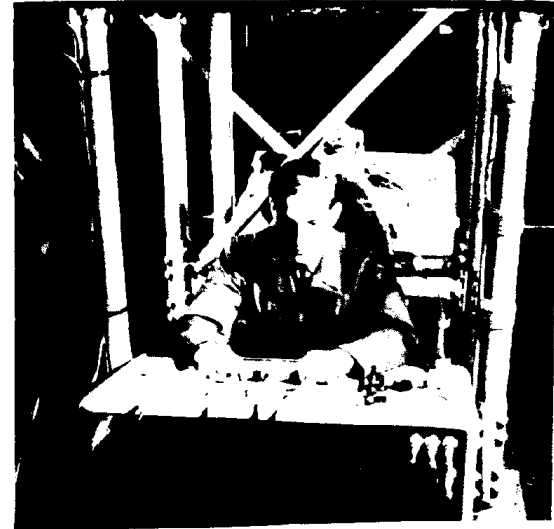
able response. (Photographs by John Halliday & Company. Similar sequences were included in a physics original research film shown by Dr. D. A. McDonald and Mr. Halliday at a *Conversazione* held lately by the Royal Society.)



22. Man-carrying centrifuge installed lately at the R.A.F. Institute of Aviation Medicine, Farnborough: the greatest acceleration which can be produced is 30 times that of gravity.



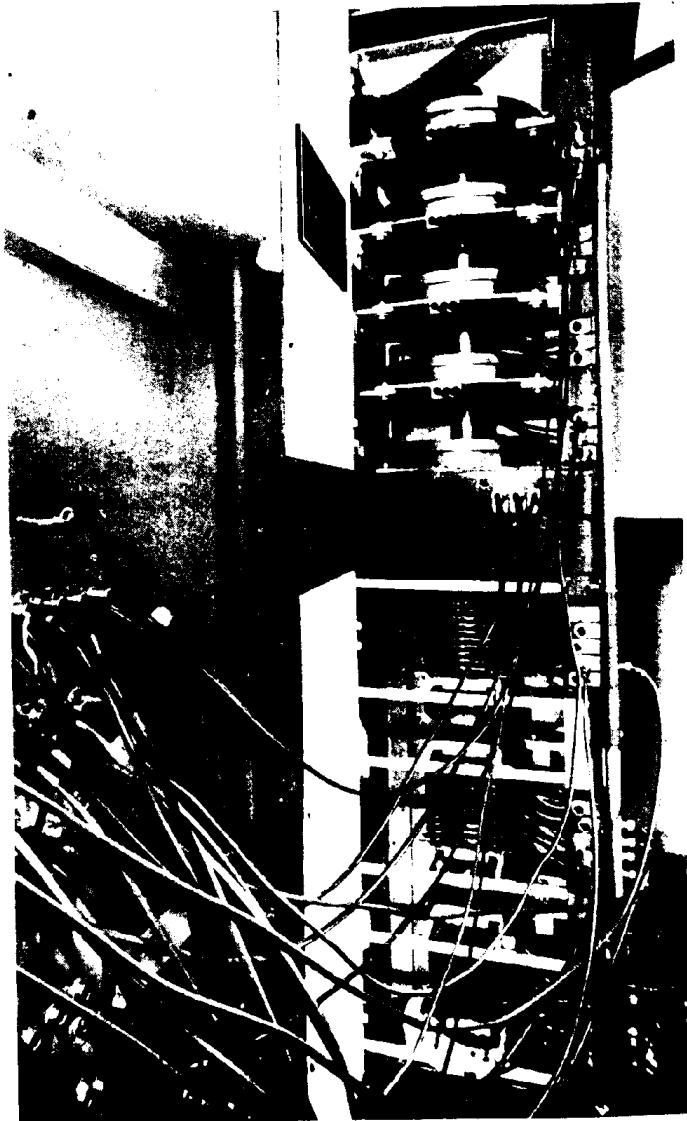
23. Man-carrying centrifuge: centre section, showing rollers and brake shoes. (Photographs on this and the following two pages are by courtesy of M. B. Wild and Company, Ltd, and the General Electric Company, Ltd.)



24. Man-carrying centrifuge: two observers sit back to back at the centre of the rotating arm.



25. Man-carrying centrifuge: a subject in one of two cars attached at distances of from 10 ft to 30 ft from the centre.



26. Man-carrying centrifuge: slipping column for medical and 'tele-talk' circuits.

to light grey. In this way, performance can be assessed for different object sizes and different object contrasts.

The results of extensive studies with this technique have enabled a code of illumination requirements for many different tasks to be drawn up in Great Britain by the Illuminating Engineering Society.* This code has been able to cover a wide variety of tasks, since once the apparent size of the detail to be recognized and the degree of contrast involved are known, the experimental data indicate the performance to be expected with various amounts of illumination. The actual level of performance aimed at in drawing up such a code is determined by economic considerations. To achieve the maximum possible performance in ordinary tasks would require illumination levels of the order of those achieved in daylight out of doors; and, since even nowadays these levels of illumination are still expensive to produce, it is generally necessary to accept a standard of performance rather less than 100 per cent of the maximum possible. This permits reduction of the illumination values by a factor of 10 or 100, since response to illumination is not linear, and performance improves only very slowly at high values of illumination. In Great Britain, a standard of 90 per cent of maximum possible performance is agreed to be an economical one at the present time. In the United States, however, where power is relatively cheaper and labour costs greater, they can afford to adopt a standard of between 90 per cent and 100 per cent of maximum possible performance. This explains why the values of illumination necessary for a given task are shown as higher in the American Illuminating Engineering Society Code than in the British Code. There is no difference between the seeing ability of the Americans and the British, but the American economic picture enables a higher level of performance to be aimed at.

That a higher possible level of visual performance in most industrial tasks provides conditions in which better productivity

*This body is the recognized centre of lighting information in Great Britain. Its members include physiologists, engineers, physicists, architects, and designers, as well as suppliers and users of light.

can be achieved is indisputable, but whether these conditions will be taken advantage of is another matter. It would seem that many factors in the running of a factory contribute to the actual level of productivity achieved. There is a well-known story of one industrial firm in the United States which tried out the effect of increasing levels of illumination on rates of production. The increases in production rate did not appear to be clearly related to the increases in illumination in every case, so a check test was made on a selected group of workers, in which the illumination was decreased instead of increased. The astonishing result was that the production rate again increased until an illumination level was reached at which the workers complained that they could no longer see to do their work. This seems to teach us that a little evidence of the management's interest in the working conditions of its staff often has a better effect than the actual easing of the tasks which results. This must be especially true when the work is not taxing the workers' visual capabilities to the limit.

QUALITY

The provision of a suitable quantity of illumination is only part of the requirement of good lighting, or even of rendering the object visible. An object is seen by reason of the contrasts arising, which may be contrasts with the background, as for instance when a needle is held up to be threaded, or internal contrasts in the task itself, as in this printed page. Shape and texture also play their part in assisting visibility, and where the contrast is not two-dimensional, as it is in reading, so-called 'shadowless' lighting is very likely to result in reduced visibility. More light coming from one side than another is usually helpful by enhancing the contrast inherent in three-dimensional objects by the production of slight shadows and by the differential illumination of surfaces at different angles to the incident light. Inset 10 and 11 illustrate the difference in visibility of set type when illuminated by 'shadowless' lighting and by lighting with a strong directional component.

One of the most important factors determining the quality

of lighting is that of visual comfort. In these days when high productivity is the aim, and any fall in quality of the work usually means loss of money, it is important to keep the workers' visual efficiency and alertness as high as possible. Uncomfortable seeing conditions are very quickly productive of tiredness and poor morale, and can completely vitiate any advantages of providing a high level of illumination.

The problem of visual discomfort has come to the fore comparatively recently as a result of increases in levels of illumination made possible by increased efficiency of light production. Many factory managements found that when they increased the illumination levels given by old lighting installations, by using a source with higher light output or by designing a new more powerful installation on the lines of the old one, the benefits they expected were not achieved. The problem arose first in the United States, the country which was the best placed economically to take the lead in raising levels of illumination. These new installations were often complained of as 'too bright' or giving 'too much light', regardless of the fact that the levels of illumination provided were only of the order of one-tenth of those which are found comfortable in daylight in the absence of sun-glare. Very little progress could be made in installation design, however, until techniques of evaluating glare discomfort were developed.

It is convenient to distinguish between two kinds of glare; disability glare and discomfort glare. Disability glare is the sort which reduces one's ability to see the task, but does not necessarily produce discomfort. Discomfort glare is the sort which produces discomfort in some degree without necessarily affecting the visibility of the task.

You can demonstrate to yourself the difference between these two forms of glare if you go near a large window and hold this book up so that the sky forms a background to it while you try to read. Unless the day is very bright, you will not find the sky itself particularly uncomfortable even to look at directly, but you will notice a remarkable improvement in visibility if you move the book so that the sky immediately around it is screened

by buildings or trees or a piece of card. This is mainly disability glare. If now, you go into a darkened room and hold the book close behind a bare lamp bulb so that the lamp is between you and the book, you will be able to read the print perfectly well looking past the lamp, but until you screen your eyes from the direct lamp light, you will be well aware of the strong discomfort glare produced. Screening your eyes from the lamp, however, will not be found to produce any very great improvement in visibility of the print.

Disability glare is much more amenable to objective study, as change in visibility threshold under various conditions of glare can be fairly readily measured. Results of detailed studies of this aspect (in particular those of W. S. Stiles in Britain) have therefore been available for some time. Discomfort glare on the other hand is an entirely subjective phenomenon, and required the development of a technique of obtaining consistent subjective judgements from a number of subjects.

HUMAN GLARE-METERS

The technique of subjective study, which has been used in Great Britain at the Building Research Station, and which has perhaps produced the most detailed data on discomfort glare, is the multiple criterion technique, developed by R. G. Hopkinson. Using this technique, the physical conditions are produced (most conveniently on a model scale) in such a way that the subject is able to regulate them by adjustment of a physical variable to produce the gamut of discomfort sensation from imperceptible to intolerable. Thus, in the studies on discomfort glare in interior lighting, on looking into a box, the subject saw the representation of a room which was apparently illuminated by bright sources of light. By a suitable optical arrangement, however, the source brightness could be varied independently of the brightness of the scene, and it was possible for a given setting of the source brightness, to make the scene itself so bright that no discomfort at all was experienced from the sources, or so dim that the sources appeared intolerably bright. The subject was given four criteria of glare discomfort, and asked to set the

scene brightness so that the discomfort reached, in his estimation, the level given. The criteria were adequately explained to each subject, and it was emphasized that the settings obtained would be entirely personal and expected to vary from individual to individual. The procedure was that the sources were set at a fixed brightness, and the scene presented to the subject at a very low brightness. At this level, the glare of the sources would be found intolerable, and the subject was instructed to raise the brightness of the scene until a point was reached where the glare was on the point of changing from intolerable to merely uncomfortable. This was defined as the criterion 'just intolerable'. Having noted this point, the subject raised the scene brightness to the point at which he considered the glare to be changing from uncomfortable to acceptable. This criterion was called 'just uncomfortable'. A similar change-over point was found where the discomfort had been reduced to 'just acceptable' and a fourth, where it had been reduced to only 'just perceptible'. The scene is shown in Inset 12, but any attempt to reproduce the glare stimuli by printing on paper is doomed to inadequacy.

In practice, it was found that the subjects were soon able, with a little experience, to set the scene brightness to the criteria given with a fair degree of confidence, and to settle down to a reasonably good level of consistency. A control set of conditions was standardized and throughout the course of the experiments the subject was 'monitored' by being given the control set at intervals. Figure 2 shows how one particular human 'meter' maintained its calibration over a long period. The standard deviation of these readings is 0.18 (log brightness units), which is a very satisfactory level on which to base a collection of psychophysical data.

There are, of course, large variations in sensitivity from individual to individual and it is necessary to decide what level of sensitivity to adopt as a design basis. The glare studies were conducted in detail with a small panel of practised subjects (to obtain consistency in readings), and covered the effect on discomfort of different source brightnesses, source areas, shapes, numbers, and positions in the field of view. The work was later

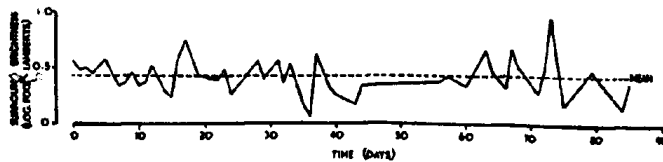


Fig. 2. Long-term consistency of subjective judgement: a typical set of observations in a control experiment. The assessment of 'surround brightness' with which given glare sources give 'just uncomfortable' glare.

extended to study the effect of grading the brightness of the source in stages from the centre of the source outwards to the surround, instead of having a uniformly bright source giving maximum contrast with the surroundings at its edges. The experiments were repeated in outline with a large number of unpractised subjects to relate the results to the reactions of as large a sample of the general population as practicable, and it was found that the skilled subjects were, if anything, more sensitive to glare than the naïve subjects. A further check of some of the conditions on full scale showed that the results obtained from the detailed studies could safely be used as a basis for lighting design, with the added safeguard that conditions designed using the average results of the skilled subjects as limits, would be likely to produce discomfort for only a small proportion of the general population.

It has been possible to derive general formulae from these studies to serve as a practical guide to the effect of many of the factors affecting glare discomfort. Workers in other countries have been carrying out similar studies using related techniques, and they also have produced formulae which have proved to be closely similar to these. Discussions will be held at the meeting of the Commission Internationale d'Éclairage (the international body for exchange of ideas and information about lighting matters), which is being held in Zürich this year, to come to an international agreement on a formula which will represent a compromise between the slightly different results of several workers.

GOOD LIGHTING

The actual results of these studies are too detailed to discuss here. The most practical benefit has been that it has been possible to draw up general rules for guidance in the design of lighting installations and fittings. To those who are interested in designing their own lighting or who have reason to suspect that their lighting may not comply with the accepted standards of comfort, two basic generalizations may be of interest. (These rules have been adopted by the Ministry of Education in framing the regulations for the guidance of designers of lighting fittings and installations in classrooms.) They are as follows:

(1) No part of any lighting fitting which is visible at an angle of less than 45° above the horizontal shall have a luminance (brightness) of more than 1,000 foot-lamberts* (average) or more than 1,500 foot-lamberts over very small areas. [The brightness of the pearl bulb of a 150-watt lamp is nearly 60 times this average value and the brightness of a fluorescent lamp 1.5 to 1.8 times. Thus, a well-designed fitting will incorporate adequately diffusing screens or louvres to ensure that the bare sources are not seen by the eye at angles of less than 45° above the horizontal.]

(2) A reasonable proportion of the light from the lamp must be directed upwards on to the ceiling. [This not only helps to eliminate the gloomy tunnel effect evident in so many industrial lighting installations, but reduces the contrast between the fitting and its background.]

Thus the fittings designer now has a framework within which fittings giving adequate visual comfort may be designed. It is up to him to evoke the right atmosphere by the use of shape, size, colouring, and sparkle, and perhaps on occasion by a well-conceived departure from the rules themselves.

*A uniformly diffusing source of luminance of one foot-lambert emits one lumen per square foot of surface. (A source emitting one candela (candle power) in all directions emits a total quantity of light of 4π lumens.)

MODERN LIGHT SOURCES

As I have already said, the aspect of lighting which attracts most public attention is the use of a new light source, particularly when it is one which gives light of an unaccustomed colour. The main reason for the rather rapid spread of the use of the discharge lamp, both with and without a fluorescent powder coating, is the increased efficiency of light production possible. The practical limit of efficiency of a lamp using a tungsten filament, heated to incandescence by an electric current, seems to have been reached by virtue of the approach of the temperatures now used to the melting point of the tungsten. The nearer the temperature approaches the melting point, the shorter the life of the lamp, and it is the inescapable inverse relation between efficiency and life which sets the limit to what is feasible with the incandescent lamp. These particular limitations do not apply in the case of a discharge lamp, in which the electrons of the gas or vapour are excited directly by an electric field instead of by heating, once the difficult technical problems in constructing a practical lamp have been solved.

Two of the ways in which discharge lamps differ from incandescent lamps have a great effect on application for the lighting of interiors. The first is that the gas or vapour, when excited under conditions of low or moderate pressure, emits the major part of its radiation in small regions of the spectrum, the number and position of which are characteristic of the atomic structure. (Unfortunately vapours cannot generally be mixed in the same discharge tube to blend their light.) The second is that an electric discharge possesses the characteristic of negative resistance. The effect of this is such that an increase of current through the discharge would bring about a reduction of the potential difference across it which would lead to further increase in current, until it reached catastrophic magnitude if a current limiting device was not included in the circuit. Such a current limiting device can be fairly simple, but it adds to the expense and complication of the installation, and hence offsets some of the advantage in efficiency.

The restriction of emitted radiation into narrow regions of the visible spectrum has until recently confined the use of discharge lamps mainly to exterior lighting. We have been familiar now for some twenty years with the blue-green light of the mercury vapour street lighting, and the yellow light (all in one very narrow spectral band) of the sodium vapour lighting. Those who have not experienced their use in interior lighting can easily imagine the peculiar effect produced when colour appearance is sacrificed to efficient production of high illumination.

The problem of how to fill in the gaps in the spectrum of radiation emitted has been almost solved in the case of the mercury discharge by the use of fluorescence. The normal low-pressure mercury discharge emits a large part of its radiation in the short wavelength ultraviolet region which is completely wasted as far as vision is concerned, and is in any case absorbed in the glass of the discharge tube. With the aid of a suitable fluorescent material, the energy of the ultra-violet light can be absorbed and re-emitted in those parts of the spectrum where it is most required. The success with which this has been achieved is witnessed by the extremely rapid spread of fluorescent lighting for nearly every purpose throughout the world in the last decade.

The modern fluorescent lamp has achieved a remarkable degree of colour correction, by relying on the use of fluorescent powders excited by the ultra-violet radiation in the discharge to produce visible light, rather than on the discharge itself; but the struggle still goes on to combine satisfactory colour rendering with high efficiency. Two separate considerations should govern the selection of the appropriate type of lamp for a given job: the colour of the light as emitted, and the colour rendering of objects illuminated by the light. These two are not the same thing; an apparently white light may be made by the mixing of two complementary colours, but the rendering of the colours of most objects by this light would be very unsatisfactory. The eye integrates the spectral distribution of the visible light it receives to perceive the final colour, and in the absence of any comparison standard it will accept a large range of very pale

colours as white, with the distinction that a tendency towards the red region is considered as a 'warm white' and towards the blue region as a 'cold white'. This fact has been recognized by the lamp manufacturers who, with the range of fluorescent powders available, are now able to produce any shade of apparently white light from very 'cold' to very 'warm'. The colours of the objects that we see, on the other hand, are governed by the relative amounts of radiation energy in the different regions of the visible spectrum which they receive from the light falling on them, and by the proportion of this incident radiation which is reflected at each wavelength. Not until after this reflection is the light integrated by the eye to perceive the colour of the object.

We are normally used to viewing colours illuminated by light which originates from daylight or incandescent bodies. The light from these sources is characterized by a smooth distribution of energy throughout the visible spectrum. It is true that the energy level rises considerably at the red end of the spectrum of light from an incandescent lamp, and at the blue end of the spectrum of daylight, but there are no sharp discontinuities or sudden changes, and the eye quickly becomes adapted to the 'warmer' or 'colder' characteristics of the source. Discharge lamps on the other hand, as I have said, emit a large proportion of their radiation in narrow regions of the spectrum, so that whatever the colours of the objects viewed, the reflected light can only fall within these particular regions. In consequence, unless the reflection characteristics of the materials happen to be such that they absorb very strongly the other colours in any case, their appearance when viewed under this lighting will be considerably changed from that to which we are accustomed. This sort of change is one to which, for some reason, the human eye is not at present able to adapt, and colour distortions caused by sharp discontinuities in spectral radiation are generally unacceptable for interior lighting where the visual environment is of any consequence.

Although the technique of using luminescent materials has developed to such an extent that the major part of the light output from fluorescent lamps is given by the continuous radiation,

which largely fills in the gaps of the visible spectrum, there are bound to be present narrow bands of the mercury radiation superimposed on the continuum, particularly in the violet-blue, green, and yellow regions, as can be seen in Figure 3. This shows the spectral energy distribution of the normal 'daylight' lamp compared with the 'de luxe warm white'. (Such a distribution is obtained by measuring the energy in spectral bands of a finite width, and any point represents the average energy in that band width about a given wavelength. The bands including the mercury vapour radiation lines are therefore only representative in area of the total energy in the band, and with a finer resolution spectrometer would be shown as much narrower and higher.) It is these bands which cause most of the disturbances in the rendering of colours. It is interesting to compare also the visual effect of the radiation from a tungsten filament lamp with that from

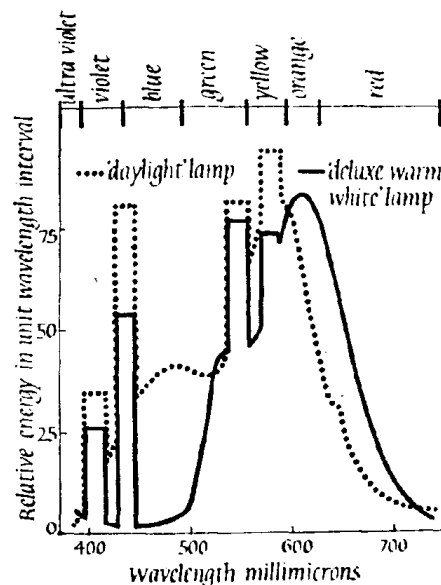


Fig. 3. Spectral energy distribution of fluorescent lamps.

a fluorescent lamp (the 'de luxe warm white') made to match this light as far as possible. In Figure 4 the spectral energy distribution curves have been multiplied by the spectral sensitivity of the average human eye. It will be seen that the two curves lie reasonably close together, although the filament lamp still predominates in the red, but there is likely to be a considerable emphasis on any light reflected in the far blue if it includes the region around 440 millimicrons, and to a lesser extent in the green at around 545 millimicrons, when coloured objects are illuminated by the fluorescent source.

Unfortunately, the more the spectral radiation is smoothed out, and built up in the orange and red regions, the less efficient

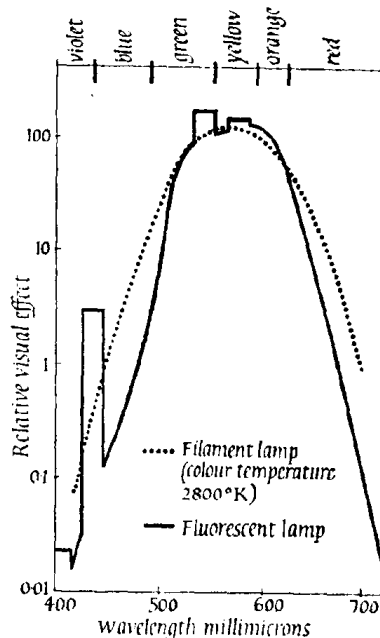


Fig. 4. Relative visual effect of light from tungsten filament and 'de luxe warm white' fluorescent lamps.

does the conversion of energy become (with fluorescent powders in present use). The solution to the problem adopted by the manufacturers at present is to make available a range of fluorescent lamp colours producing 'warm' and 'cold' 'white' light from which the user can make his choice of greater or lesser efficiency, according to how seriously some distortion of colours will affect the work to be done or the effect to be produced. None of these lamps can give colour rendering which matches perfectly the rendering by a given kind of daylight or by filament lamps, but 'northlight' or 'colour matching' tubes are available, which produce a very 'cold' light similar in appearance to that from a north sky which is frequently used as a standard for colour matching. These are found suitable for many situations where good colour rendering is required, but the light they give is not very 'cheerful', and they are much less efficient than the 'daylight' type. To produce light of an appearance matching that of a tungsten filament lamp with reasonable colour rendering at some sacrifice in efficiency, a 'de luxe warm white' lamp has recently been introduced, while for a more deliberately emphasized 'warm' effect, colours such as 'peach' and 'mellow' are available.

The choice of 'cold' or 'warm' lighting will be based mainly on the atmosphere that is required to be created by the lighting. Where the light is to be blended with daylight, or in factory shops where hot processes are carried on, and for general use in hot climates, the 'cold' lamps may help to produce a more satisfactory atmosphere, whereas for social purposes or in cold places the 'warm' colours will usually be preferable. It is interesting, however, to note the work of A. A. Kruithof in Holland and A. Wald in Germany, who have established that for general purposes a rather high level of illumination is required from 'cold' lighting if it is not to appear inadequate and gloomy. This is possibly because of our association of the cool lighting with daylight, and we get an impression of inadequacy if the levels are far below those received indoors with good daylighting, but it is a rather difficult subject to study with any degree of precision. It is, however, generally accepted that for

lighting to levels of illumination provided by filament lighting, for general purposes, a 'warm' or 'natural' type of tube is usually to be preferred.

Fluorescent lamps are at present available in two distinct forms. Both use a discharge in mercury vapour at low pressure as the primary source of radiation, but in one form, for general use, the discharge is fed directly from the (200-250V) supply mains, in series with the current limiting device, and in the other form it is fed through a transformer supplying high voltage, but with sufficient impedance to limit the current to the correct value. The mains voltage lamps are restricted in length to not more than 8 feet, and it is necessary to heat the electrodes to produce sufficient emission to start the discharge. This is accomplished by the use of an automatic starter switch whose function is to apply the mains voltage to the tube when the electrodes have heated sufficiently, or by connecting both heaters to a specially wound transformer (the 'instant start' circuit) which dispenses with the starter switch. This type of lamp is termed a 'hot cathode' lamp. The other type, the 'cold cathode' lamp, is not restricted in length, which only depends on the voltage provided by the transformer, and it strikes instantly without the necessity for heating the electrodes. An installation using cold cathode lamps will usually be more expensive in first cost than one using hot cathode lamps, but there is the great advantage that the life of the tubes is of the order of 15,000 hours, or three times as long as that of the hot cathode lamps. This makes them particularly appropriate for use in situations where replacement of lamps is unusually expensive or inconvenient, owing to disturbance caused to a production area or to the necessity of special scaffolding to reach great mounting height.

The low-pressure mercury vapour fluorescent lamps which I have been discussing have the feature that they are large in physical dimensions, although they are only of relatively low wattage. While this is an advantage for many types of lighting, it is sometimes a disadvantage for large industrial plants with a high roof, where it is often more convenient to use a small

number of high-powered light sources in fittings which concentrate the light on to the working plane, without letting them illuminate the upper part of the building more than is required for producing a visually comfortable working environment. It is in these situations that the uncorrected high-pressure mercury vapour lamp, on account of its high efficiency and availability in powers up to 2,500 watts, has been used in interiors, with rather unpleasant effects on the appearance of human beings and many of the objects around them. A blend of mercury light with tungsten filament light improves the colour rendering considerably, but of course lowers the overall efficiency. Another development of photoluminescence has been of great help in this field, providing a powder which, when activated by the long-wave ultra-violet band in the high-pressure mercury spectrum, fluoresces with a strong red colour to add its quota of light in the region where the main deficiency of the visible spectrum of this type of discharge lies. This provides a very efficient high-powered source of light of pleasant colour for use where accurate colour rendering is not vital.

FLICKER

One of the characteristics of vision which has had to be re-examined in recent years is that of the perception of flicker, and the discomfort caused by visible flicker. This re-examination was the consequence of a very small but persistent trickle of complaints from among users of fluorescent lighting that this had been found an unsatisfactory type of lighting in some instances. It seemed on consideration of many cases that the chief causes of complaints were the poor design of the installation leading to visual discomfort and glare, and the use of unsatisfactory lamp colours. These two factors have already been discussed above. There is, however, the third factor of flicker which was occasionally specifically mentioned as a cause of complaint.

Discharge lamps, whether plain or fluorescent, are usually run from alternating current supplies which are standardized at a frequency of 50 cycles per second in Great Britain and most of Europe, and at 60 cycles per second in the United States. The

light output from the discharge, which does not have the thermal inertia of the incandescent filament, follows closely the current input and falls to zero every half cycle, producing a light fluctuation at 100 cycles per second. Some 'inertia' is introduced into the light emission by the use of the fluorescent powder coating, but it is not usually possible to increase the minimum light output during the period of zero current to a value greater than 50 per cent of the maximum. The most familiar consequence of the fluctuating nature of the light is the 'stroboscopic' effect which is seen when objects are moved under lighting of this kind. A bright object moved across a dark background will appear to move in a series of jerks, and a rotating one may appear to rotate at a speed different from its real speed, or even to be stationary or rotating in the opposite direction, according to the relation between the real speed and the frequency of the light fluctuations. This stroboscopic effect was known before the introduction of discharge lamps, and the danger of their use to illuminate machinery without special precautions in the installation has always been stressed. It was not, however, until the advent of the highly efficient fluorescent lamp, encouraging high levels of illumination to be provided on the work, that the problem of flicker seen on stationary surfaces demanded attention.

The original work on which prediction was made as to the visibility of flicker from discharge lamps was that of R. J. Lythgoe and K. Tansley in 1926, who carried out observations using small fields of view. From the results of this work, it had been assumed that any variation at a frequency as high as 100 cycles per second would not be perceptible to anyone, but it appeared that nobody had studied the effect of stimulating the whole of the retina, as would occur in a lighting installation. It was, therefore, decided that the Building Research Station should try to obtain data under these conditions on the average sensitivity of flicker and on the variations in sensitivity from subject to subject, and from day to day for the same subject.

This investigation was carried out by a subjective study similar to that described for studying the reactions to glare. One of the most important results of this investigation was the dis-

covery that the variation in sensitivity to flicker occurring over the whole field of vision between one person and another is extremely great, apart from the large variation which may occur for the same subject from day to day. The histograms of Figure 5 show the spread of results for 20 observers under a fixed set of conditions in terms of frequency of fluctuation required for four different degrees of flicker sensation. It will be seen that flicker which is found to be quite uncomfortable to one man may not be perceptible to another, so that if only one worker complains of flicker in an office where 50 are working under fluorescent lighting, there is no justification for telling him he is imagining it. He must either be provided with conditions less

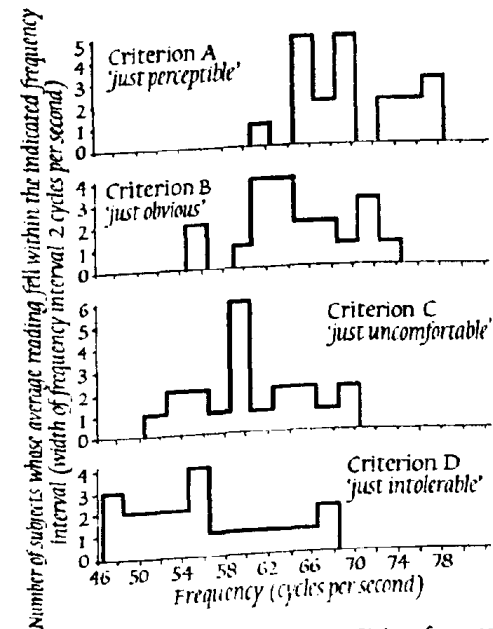


Fig. 5. Distribution of average values of flicker frequency for 20 subjects for four criteria of flicker sensation (full scale experiments).

likely to cause perception of flicker or be told to work elsewhere.

The investigation has shown how the perception of flicker is affected by such factors as brightness of the field of view and wave-form of the fluctuations. The greater the brightness, of course, and the greater the 'severity' of the fluctuations, the greater is the possibility of the perception of flicker at a given frequency. It has been shown that the possibility of perception of flicker from fluorescent lamps functioning properly is extremely low, except perhaps under conditions of very high field brightness, such as may occur in, for instance, a drawing-office which is lighted to a very high level of illumination to enable fine work to be done without any local lighting. This possibility can be reduced by using three lamps close together, fed from different phases of a three-phase supply. In this way, the fluctuations from the lamps can be made to overlap, as the alternating currents in the three phases will be 120° out of phase. If a three-phase supply is not available, some measure of smoothing of the fluctuations can be obtained by running two lamps with their currents displaced in phase, by using a condenser as part of the limiting impedance for one lamp, and a reactance for the other. The current in the lamp using the condenser ballast will be in advance of that in the lamp using the reactance, and the light wave-forms will again overlap. Figure 6 shows the effect on light-time wave-form of changing from a single lamp, first to this 'lead-lag' circuit using two lamps, and finally to the three-phase arrangement using three lamps. The superior smoothing of the three-phase arrangement is very obvious.

It has been found that the light output wave-form of some fluorescent lamps is not symmetrical on both halves of the current cycle. This has the effect of introducing a small component of 50 cycles per second fluctuation which is much more likely to be seen as perceptible flicker. It is nearly always seen near the electrodes of a lamp, which is why the ends should always be screened for two or three inches along the tube, but the effect does not usually extend very far along the tube. When the whole tube is affected, because perhaps of asymmetry of electrode emission, it has been found that very small amounts of 50-cycle compon-

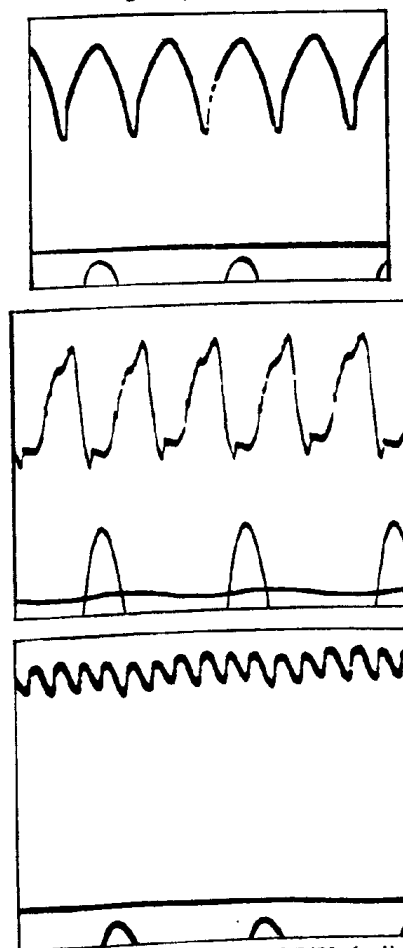


Fig. 6. Light output wave-forms of MCF/U daylight-type lamps, recorded on cathode-ray oscilloscope: (a) single 80-watt lamp; (b) two 40-watt lamps in 'lead-lag' circuit; (c) three 80-watt lamps fed from three different phases of a 3-phase supply. The supply frequency is 50 cycles in each case, and the supply peaks are shown at the bottom of the records.

ent cause visible flicker at moderately high levels of field brightness. The effect is not necessarily associated with lamps near the end of their life, and the evidence seems to be that most of the complaints about the flickering of fluorescent lighting (where not referring to the stroboscopic effect) arise from the presence of lamps exhibiting some proportion of 50-cycle component rather than from the inherent 100-cycle fluctuation of the light. It would seem that the use of twin-tube fittings (say, two 40-watt lamps) instead of single tube (one 80-watt) would help to reduce trouble from this source, but fundamentally it is a problem of manufacture.

COLOUR AND WORK

The period since the war has seen a great increase in the use of colour in interior decoration, particularly in relation to the visual environment in factories and offices. There is no doubt that the competition for labour has played a large part in persuading factory managements to improve the environment of their workers, but many have been surprised to discover what difference good lighting and a proper colour scheme have made to the standards of workmanship and morale. A gloomy and dull factory is bound to be a dirty place, and the 'housekeeping' and quality of work is almost certain to suffer. A well-decorated work-place which is not too difficult to keep clean and which is well lighted so that dirt and rubbish can quickly be noticed, will, in fact, be kept clean if it is rewarding to do so. On the other hand, attempts to turn a busy production plant into a kind of sterilized dairy shed will be certain to cause disappointment. A scheme must take into account the character of the work-place and the processes involved. A scheme which is suitable for a food factory or a textile finishing shop will usually not be appropriate to an engineering works. Due consideration must be given to the inevitable spread on to machines of such materials as cutting oil, graphite, chalk, pigments, etc., which can ruin a 'theoretical' colour scheme in minutes – and make the task of keeping the intended character of the place Herculean.

The details of choice of a colour scheme must always lie with the designer's artistic sensitivity, but the development of an approach to the design of the brightness relations in the field of view has enabled the colour scheme designer to base his choice on certain fundamental principles. The starting point is always the work itself. This should be the brightest and most attracting point of the local field of view. Colours of the machines and the walls on which the eyes can rest while relaxing from the task will generally be kept darker in tone than the work point, and will not contrast too violently with it. Colours on machines and surroundings will be generally not of high saturation, and chosen so that inevitable exudations from the processes do not immediately make the scheme worthless. On the other hand, surfaces seen against or close to windows, roof lights or artificial lights will be kept as light as possible, even white. Areas of strong colour to give the main character to the work-place will probably be usually kept to surfaces which are normally out of the field of view of the man concentrating on the work. The actual hues and strengths of these colours will depend on the character which the designer wishes to produce for the space, and the nature of the work being carried on. For instance, in a factory area where the processes are largely automatic and repetitive, he may use bold and saturated colours, whereas where the work demands long periods of visual concentration he will probably use less stimulating colours and patterns. Other considerations such as the basic colours of the materials being processed, and the desirability of affecting the feeling of warmth or coolness in the area, may determine the actual hues used.

A systematic arrangement of colours is of great value to architects and designers working with colour. The problem of the systematization of colour and the attempts made to solve it have been discussed by C. L. Boltz in *Science News* 22. It has been found that the system most satisfactorily linking physical and subjective variables for use with a technique based on fundamental principles, is the Munsell System. (See pp. 49–52 of Boltz's article.) This is based on hue (whereabouts in the spectrum), 'chroma' (the attribute of a colour by which a judgement

is made of the amount of pure chromatic colour present irrespective of brightness) and 'value' (the lightness or appearance of reflecting a greater or smaller proportion of the total incident light). With this system, the designer working out a colour scheme can have before him a series of scales from which he can choose which value and chroma of a particular hue is the most appropriate for the visual requirement, or which hue and chroma of a particular value gives the right colour balance and atmosphere. These colour scales can be viewed under the type of source chosen for lighting the area. This is a vital precaution; many a disappointment has been incurred by not allowing for the change in appearance of a colour scheme when the illumination is changed from daylight to (say) fluorescent lighting.

FUTURE RESEARCH

The study of the influence of lighting on the work and well-being of the individual has had vast new fields of research opened to it by the development of techniques of subjective study, and these in turn have demanded the development of more detailed methods of objective study of the physical characteristics of lighting installations.

Nevertheless, it is probably more difficult for anyone not working in the particular field to envisage the direction of future research in lighting than in many other fields. It is taken for granted that much work is needed on problems of disease, development of materials, fundamentals of physics and physiology and even on the processes of vision, but it is much less obvious to anyone who has not contemplated the problems involved what facts about lighting still remain to be discovered, apart, of course, from that eternal dream – the perfect lamp! To those working in the field of lighting, however, there are many questions which pose themselves, the answers to which, we are convinced, would help us considerably towards achieving the 'perfect' lighting installation.

One such question is whether many of the difficulties of present fluorescent lighting might not be eliminated at their source

by dispensing with the gas discharge tube. Why produce unwanted kinds of light only later to convert or 'correct' them by fluorescence? Why not apply electrical energy to stimulate the emission of light directly? With this idea, there has been in recent years a renewal of interest in the phenomenon of electroluminescence. Briefly the arrangement for using it as a light source is to incorporate the luminescent phosphor in the dielectric of a simple condenser with transparent plates. The transparent plates consist of glass with an extremely thin conducting metallic film deposited on the surface. The amount of light emitted by the phosphor depends on the voltage applied to the plates of the condenser and its frequency. So far, however, it has not been possible to obtain an efficiency of light production greater than about 1/5 to 1/10 of that obtainable from the common fluorescent lamps. The mechanism of the phenomenon is, moreover, obscure; and, at the best, a great deal of fundamental and development work will be necessary before a revolution in techniques of lighting could be produced.

The illuminating engineer trying to plan a lighting technique from fundamental principles comes up against gaps, too, in knowledge of basic visual factors. One of these is the phenomena of visual fatigue. What is visual fatigue? Does it really exist? Are there objective signs or measures? So far we have only subjective impressions to go on, but it is not disputed that a day's work under unsatisfactory lighting conditions is more exhausting than a day's work under good, well-designed lighting conditions. Experimental psychologists have made a start at studying this phenomenon, and a reasonably promising technique is being applied in lighting research.

An allied phenomenon is that of the mechanism of attention. It is known qualitatively that the eye is instinctively drawn to the brightest part of the field of view, but the extent of distraction caused by such effects as reflected glare or unsatisfactory brightness distribution is unknown.

The development of light sources of improved colour depends on a knowledge of the processes governing apparent colour rendering. Investigation of these processes is notoriously difficult,

but some fundamental work on permissible change in spectral characteristics for just noticeable colour change of pictures has recently been described by B. H. Crawford. This should provide data which will have considerable direct bearing on tolerances in spectral distribution of light sources.

To be able to predict the effectiveness of a design, it is necessary first to be able to calculate the complete brightness pattern in the interior, and secondly to know the relationships between measured brightness (luminance) and apparent brightness (luminosity), and much work is going on at the present time to study methods of determining these factors. The brightness of any one surface in a room is dependent on integration of the light received from all the other surfaces and objects, which in turn depend for their brightness on the light received from all the other surfaces. This problem can be treated mathematically, as it has been in America, but the procedure is so laborious that it is very rarely used for design purposes, so that attempts are being made in Britain to arrive at an understanding of the problem, using model rooms and installations. It is hoped to arrive at simplified (although probably somewhat less accurate) treatments for special cases of day and artificial lighting. The apparent brightness of objects or surfaces as seen by the eye depends on the level of brightness to which it is adapted by virtue of the surrounding field of vision. The study of the relationship of these factors and its application to general lighting phenomena was described in Britain before the last war. More recently J. M. Waldram has shown how the method can be applied to a detailed study of particular interiors. A study of the subjective effects of colour relations is another important part of this work.

It has been pointed out that completely satisfying lighting cannot be achieved by a static arrangement of brightness patterns, but that continual slow variations should be introduced, such as occur during daylight with moving clouds. The changes should not, of course, be crudely obvious, but rather imperceptible changes in atmosphere. Some experimental installations have already been put into effect to try these ideas, and it is

evident that the dynamics of lighting offers another large field for study.

I hope it will be evident from what I have said that lighting as a science draws upon many other different sciences for the body of knowledge which stands behind the technique of providing the best possible lighting for each of the many varied tasks and activities which are part of today's life. The best that is possible is sometimes limited by economics, and sometimes by gaps in our knowledge, but work of the sort that I have indicated is being carried out in many different countries with the aim of reducing the effects of both these limitations.

ACKNOWLEDGEMENTS

Figures 2, 4, 5, 6, and 7, and the photograph reproduced as Inset 12 are Crown copyright, and are reproduced by permission of the Director of Building Research. For Figure 1, I am indebted to Mr H. C. Weston, and for Figure 3 to Thorn Electrical Industries Ltd.

FURTHER READING

- Weston, H. C. *Light Sight and Efficiency*. London, H. K. Lewis, 1949.
- Hopkinson, R. G. 'The Multiple Criterion Technique of Subjective Appraisal'. *Quarterly Journal of Psychology*, 2, 124 (Part III, 1950).
- Comptes Rendus: Commission Internationale de l'Eclairage 1951* (also 1955, to be published). Central Bureau of the C.I.E., 1860 Broadway, New York. These proceedings include (in English, French, and German) the secretariat reports of progress in all member countries in nearly 40 different branches of lighting, as well as individual papers, describing new work in detail, and resolutions on standardization.

THE AUTOMATIC FACTORY

H. D. TURNER

BEFORE 1700, the only available sources of power were water-wheels, windmills, and the muscles of animals and human beings. By the end of the century James Watt had modified Newcomen's original pumping engine, Trevithick had invented the high-pressure steam engine which was to be the basis of steam transport, and the first industrial revolution was under way. Coal and iron became the basis of Britain's new industrial strength. Steam power, which made the old cottage industries obsolete, led to the development of the factories. The impact of this industrial revolution on an essentially agrarian society has left in its wake economic problems, labour difficulties, restrictive practices, and general social unrest; it has also produced a highly organized industrial society in which large numbers of people are engaged in monotonous repetitive work in mass-production factories. Recent developments in electronic equipment, however, now make it possible to carry out these repetitive tasks by electronic devices and thus to produce completely automatic mass-production factories. In these, the control and coordination of the activities of individual machines will be accomplished by electronic computing machines and such factories will replace those in which such control and coordination is, as at the moment, the concern of human operatives.

We can see how this may be done by considering a typical mass-production factory through which passes a continuous flow of articles in course of manufacture. During its passage through the factory, each article is operated on by various fabricating machines and transferred, at appropriate times, from one machine to the next. The machines are tended, and articles transferred from machine to machine, by factory operatives; and the organization of the smooth flow of articles through the

factory is also a human responsibility, as is the preparation and computation of appropriate written records of factory and machine activity. These, then, are the functions which, in the automatic factory, will be performed by electronic equipment; and an understanding of how this may be achieved requires detailed analysis of the structure and functioning of machines.

We may conveniently define a machine as a device which is used to direct the application of power so as to produce some desired effect. This may be, for instance, the movement of an object, the production of subsidiary effects which recur at definite intervals of time, the prevention of fluctuations of the temperature within an enclosure, or the removal of metal from a given piece; and the corresponding machines would be the crowbar or lever, the clock, the thermostat, and the lathe. We see that this definition is very wide in scope, and that it obviously includes fabricating machines of varying complexity as well as machines for transporting objects from one place to another.

The structural complexity of a machine is determined by the complexity of operations which it is required to perform; and we may suppose that the former must always be greater than or equal to the latter. In practice, we produce complex machines by assembling relatively simple machines in definite functional arrays; and the coordination of the functions of these single machines is arranged by the information 'built-in' to the machine during its construction. This building-in of information amounts to giving instructions to the machine about its future behaviour under different possible circumstances.

AUTOMATIC MACHINES

We may also imagine complex machines in which the future operation of a block, or blocks, of simple machines depends in various predetermined ways on the actual operations of other aggregates of simpler machines. Such complex machines are, therefore, automatic in operation, and the predetermination of function is achieved in this case by the mechanism for information transfer which is incorporated when the machine is constructed. This may be illustrated by the development of the early

steam engines. Newcomen's original pumping engine was manually operated; when the piston reached the bottom of the cylinder, steam was admitted by operation of the appropriate valve; and when the piston reached the top of its stroke the steam was shut off and cooling water sprayed in to condense the steam in the cylinder. The piston was then forced down by external air pressure and when it reached the bottom of its stroke the whole cycle was repeated. It was soon realized, however, that the movements of the piston could be used to operate the appropriate valves; and so the engine became an automatic machine.

Automatic machines, therefore, are defined as those whose future activity depends in some predetermined way on their past activity; they will repeat indefinitely a particular cycle of events; and this cycle may be changed by altering the manner in which the predetermination is achieved. The necessary features for this type of function are that the machine should possess parts which respond in a definite adjustable way to certain selected internal or external changes; and that devices should be incorporated which produce effective action in response to these.

We may further consider machines which are formed from simpler machines by arranging these in an inflexible manner. Such machines may be automatic, or manually operated, but in both cases they will be capable of performing one general function only. It is possible, however, to arrange that the mode of assembly and operating times of the simpler machines may be varied, either automatically or manually; and this gives a corresponding increase in flexibility, particularly since in each case the simpler machines may be selected from a number of similar devices.

This may be illustrated by considering certain complex machines which are already in use in the motor-car industry; and this is a particularly useful illustration since these machines represent an interim stage in the development of completely automatic production units for the manufacture of complex products.

As early as 1930, the A. O. Smith Corporation of Milwaukee, Wisconsin, were making automobile frames, in a completely

automatic plant, at the rate of 10,000 a day. The plant consisted of nine separate units, linked by automatic materials-handling devices. These units inspected incoming raw steel strips (rejecting those which did not meet required standards), cleaned, bent, and drilled the acceptable strips, assembled the frames, inserted and drove home the rivets, and, after exposing the completed frames for inspection, washed, dried, and painted them, subsequently dispatching them to store. Changing from one type of frame to another involves changing the tools and handling equipment in each unit, and this is done by human operatives. We see, therefore, that such machines are automatic in operation, and that flexibility is achieved by manual alteration of the simpler machines which collectively comprise the complete production machine; and that this latter consists of automatic fabricating machines and materials-handling devices. The use of this technique is economically sound because the extremely high output of such machines more than compensates for their high cost.

Transfer Machines

Similar developments in the production of components for motor cars have been recently described by the Ford Motor Company in America and by the Austin Motor Company, Ltd, in Britain. In the production of cylinder blocks, a number of cutting, grinding, and drilling operations must be performed on the rough casting. This is traditionally accomplished by separate electrically operated semi-automatic machines, each of which, under human supervision, performs certain operations; the partly finished casting is then transferred to the next machine and ultimately emerges as a finished block. These firms, however, have now installed what are called 'transfer machines'. These are in effect composite machines which are assemblies of separate simpler machines. These latter are placed at definite points, called 'stations', along an automatic conveyor system. At each station, some particular operation (drilling, grinding, etc.) takes place; when each operation has been completed appropriate switches are tripped and each block is moved for-

ward, or 'indexed', to the next station. There are frequent intermediate stations where automatic gauges inspect the work which has been done, rejecting those blocks which have been imperfectly treated and allowing the others to pass to the next fabricating station. Elaborate arrangements are made for the cleaning of worked parts, the removal of waste material, and for the indication of defective parts of the plant. Precise estimates are made of the expected lives of the various tools; in some cases the machines count the number of articles fabricated and indicate when the expected life of the appropriate tool has been reached; this tool is then changed, and the machine continues its operation. As each article is indexed to a station it is located exactly, by locating pins, for the appropriate machining operation; and this will not start until the appropriate electrical contacts have been made at all the stations. The machine may be altered so as to produce a different product by changing the tool assemblies and the equipment which carries the article in course of manufacture. The transfer machine is therefore a complex automatic machine formed from a relatively inflexible assembly of subsidiary machines, with automatic handling devices which pass the article from one subsidiary machine to the next; and the timing of the cycle of processes is governed by electrical signals generated within the machine by its own activity.

It is thus possible to produce composite machines which will carry out large numbers of separate fabricating operations without specific human intervention. For instance, at the Cleveland Engine Plant of the Ford Motor Company, a series of 42 transfer machines has been linked by automatic materials-handling devices; and the fabricated component (a 180-pound cylinder block) has 530 cutting and drilling operations performed on it during its course through the complete set of machines without its being touched by human hand. Inset 1 and 2 show parts of one of these transfer machines and Inset 3 gives an unusual view of a materials-handling device which rotates the cylinder block as it passes between two transfer machines. Inset 4 shows a transfer machine built by James Archdale and Company, Ltd, and installed in a modern automobile factory. In this, the

various stations and the conveyor system can be clearly seen.

The costs of installing and operating transfer machines compare very favourably with those for a conventional set of standard machines. Mr F. Griffiths, Chief Project Planning Engineer at the Austin Motor Company, Ltd, has recently given an interesting comparison between a 13-station transfer machine and the orthodox machines required to perform the same fabricating operations on a cylinder block. These operations are: (1) drilling all the holes in the front and rear faces, but excluding the oil gallery holes; (2) core drilling, semi-finishing, and finish boring of camshaft and crankshaft bores; (3) facing, opening out, and reaming oil pump bores; (4) milling the cam clearance around the oil pump bosses; and (5) facing the centre bearing crank thrust face. The capital cost of the transfer machine is £25,903 compared with £30,850 for a set of standard machines. The transfer machine occupies 390 square feet of floor space compared with 500 square feet for the standard machines; the costs per machine operating hour are respectively £3 13s 3d and £4 8s 9d; and output is increased from 2,500 to 3,000 cylinder blocks with a total saving of £123 per week. (Labour costs are 11s as against £2 17s 2d per hour.)

In these complex machines, coordination is effected by electronic control equipment; and flexibility is achieved by manually altering the individual machines (which are selected from certain standard machines) and the control equipment (which is merely apparatus for information transfer within the machine together with appropriate 'effector devices').

Universal Machines

If the simple subsidiary machines which together form a complex machine are of a very general type, and if provision is made for every reasonable interconnexion to take place, then we shall produce what we might term a 'universal machine'. By providing such a machine with the requisite instructions for assembling its constituent machines and timing their operations, it is obviously possible to duplicate the activity of any rigid assembly of such constituents. We can store the information necessary to regulate

the collective activity of various component machines in the form of pieces of material of appropriate length, shape, and contour; and with suitable ancillary equipment, these will govern the activity of quite complex machines. But changing the information content of the machine is often a difficult matter in these circumstances and, particularly where the cams, etc., are located in inaccessible parts of the machine, it often involves considerable labour. Once it is realized that such pieces of material merely store operating instructions in a mechanically accessible form it becomes possible, in many cases, to produce universal machines by using other ways of storing, transmitting, and using such information. We may, for instance, store information in the form of electrical signals by using magnetic tape; or we may generate patterns of electrical signals by punching holes in tape, and passing this through a photo-electric tape reader. These electrical signals, suitably amplified, may be passed to electro-mechanical effector devices and thus used to coordinate the operations of subsidiary machines within a more complex one. This obviously produces a very flexible system of machine control, although the extent to which it is employed will depend very much on the type of machine and product with which we are concerned. It is probable that in the automobile industry, even the completely automatic factory will include large numbers of transfer machines which are very largely self-regulating. In these, outside intervention will probably be limited to altering the timing of various operations; to adjusting the relative outputs from duplicated stations (this may be necessary because a particular process takes longer than the transfer machine cycle); and perhaps to automatically changing one set of tools for another at each station in case the tools are worn or the machine is to be set for a different product. On the other hand, it is often convenient to be able to make considerable alterations to the structure and operation of complex machines by appropriate electrical signals; some of these machines will be considered in the next section.

CONTROL SYSTEMS

A machine has been defined as a device which is used to direct the use of power so as to produce some desired effect. During the operation of the machine, therefore, this effect will gradually be achieved; there will, at any instant, be a difference between achieved and desired effect; and this difference will be zero at the precise instant that the objective is attained.

In the use of machines by human operatives, this difference is used to control the future operation of the machine; and indeed this mode of operation is fundamental to any control system. We may perhaps illustrate this by reference to two human processes. When I pick my pen from my desk, I do not smash my fingers against the wood, no matter how rapidly I move my hand at the beginning of the operation. In fact matters are arranged so that the downward velocity of the hand is zero at the precise instant that the fingers come into contact with the pen; and this is done by continuously observing the distance between pen and fingers and using the measure thus produced to control the downward velocity of the hand. This observation may be accomplished visually, or by the interpretation of nerve impulse patterns derived from nerve endings in the joints and muscles of hand, arm, and shoulders. But however the information is derived, the functional arrangement is always the same: there is a 'receptor', i.e. some method of measuring achievement at any time; there is a 'comparison device', i.e. some method of comparing this with an appropriate measure which represents completion; and an 'effector' by which the result of the comparison is used to regulate future activity. We may represent this arrangement diagrammatically as in Figure 1; the mode of operation is designated as 'negative feedback' because the system functions so as to reduce its activity; and we may note that in the case quoted above, elaborate subsidiary negative feedback systems are employed which prevent the hand from overrunning, and hence oscillating about, the pen.

A mechanic reducing the size of a piece of metal by filing exhibits the same functional pattern. We may suppose that it is

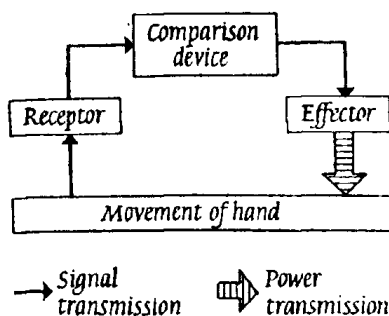


Fig. 1. Diagrammatic representation of the simplified control system involved in picking up a pen from a table. The receptor continuously produces signals which indicate the position of the hand. These are compared with signals which indicate its desired position; the difference is used to actuate the effector device which consists of the various muscles involved in the hand movement; and this difference and therefore the hand movement will both be zero at the precise moment when the hand touches the pen.

desired to remove metal until the dimensions of the residue agree with those given in a drawing. Metal is removed by filing for a given time which depends on the experience of the operative. He then measures the appropriate dimension of the work and compares this with the required value. The difference between these two measurements determines his future activity, and by repeating the process he gradually approaches, and finally attains, the desired result. We may represent this process diagrammatically as in Figure 2. It is not continuous, as in the other case cited, and this is brought out in the diagram. On the other hand, measurement, transmission of information, comparison, and effective action are all ultimately accomplished by patterns of nerve impulses, as in the first case.

We can, however, produce machines which exhibit these characteristic features of control systems. One of the earliest of these was the speed governor for steam engines, invented by James Watt, and shown in Inset 5. The position of the spheres

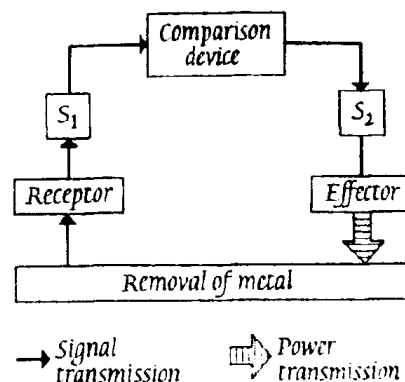


Fig. 2. Diagrammatic representation of activities involved in filing a piece of metal to a given shape. The receptor produces signals which indicate the amount of metal removed. These are compared with signals which indicate the required amount; and the difference is used to control the further removal of metal. S_1 and S_2 are switches, which indicate that the process is not continuous; in general, S_1 is open when S_2 is closed and vice versa.

depends on the speed of rotation of the output shaft; movement of the spheres, by a suitable arrangement of rods, controls the admission of steam to the cylinders; and, by suitable connexions which ensure negative feedback, the output shaft speed is stabilized at a value which depends on the position of the valves and push rods at that speed. Figure 3 is a diagrammatic representation of the device, and is exactly analogous to Figure 1. We may note also that the controlled speed may be varied by altering the setting of the valves for a particular position of the push rods. This amounts to altering the value of the comparison standard; and this is the general method of altering the level at which any controlled system operates.

While the comparison standard remains constant, the machine represented in Figure 3 is a closed-circuit machine; it operates at a level determined by this standard (which, in the context of the machine, carries information about its desired operation); and

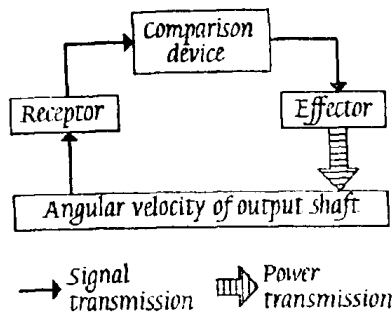


Fig. 3. Diagrammatic representation of the operation of the Watt speed governor for steam engines. Reception and comparison of signals and operation of the effector are all accomplished mechanically. Appropriate mechanical couplings ensure negative feedback.

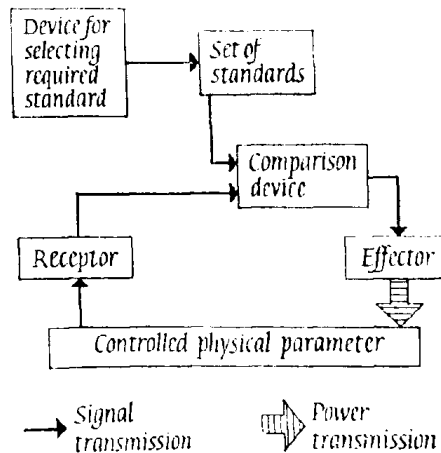


Fig. 4. Diagrammatic representation of a control system which can be operated at different levels. These are determined by the particular reference standard selected. Operation at each level is stabilized by negative feedback.

it is constrained to do so by negative feedback. We may, however, cause the comparison standard to vary; thus machines may be made to operate in accordance with external circumstances; and this operation may or may not be stabilized internally by negative feedback. The two possibilities are shown diagrammatically in Figures 4 and 5. The first of these possibilities is shown more clearly in Figures 6 and 7, which represent a thermostat and a position control servo-mechanical system respectively.

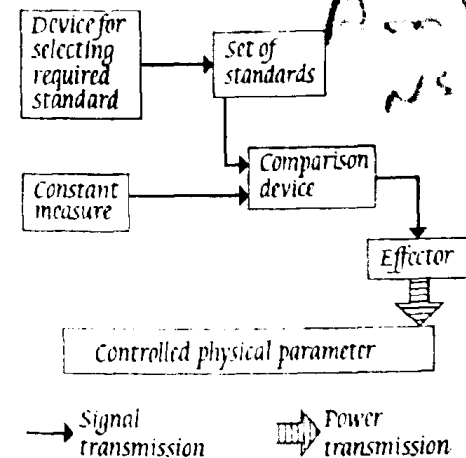
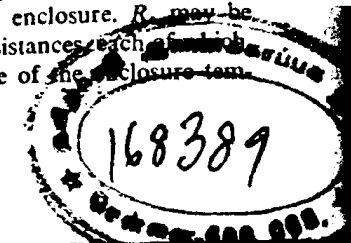


Fig. 5. Diagrammatic representation of an open-circuit control system. Here there is no stabilizing negative feedback. As in Figure 4, the value of the physical parameter is determined by the reference standard selected.

In Figure 6, the receptor is a platinum wire, the resistance of which depends on its temperature. This resistance is compared with a standard resistance R_1 ; changing R_1 , therefore, changes the controlled temperature within the enclosure. R_1 may be selected from a number of standard resistances each of which therefore, represents a particular value of the enclosure temperature.



perature. In the equipment shown in Figure 7, control is exercised by a combination of mechanical and electrical signals. The object of control is in this case to maintain a load-bearing shaft in an angular position which has been preselected. The angular position, θ_L , of the load is measured by appropriate gears and compared with the position of the handwheel, θ_H , by a mechanical differential. The difference $\theta_H - \theta_L$ is converted into an electrical voltage V_D by the rheostat so that V_D varies as $\theta_H - \theta_L$. V_D is amplified by a thermionic amplifier and thus operates the servo-motor which moves the load. The comparison standard is therefore θ_H which is manually variable. It can easily

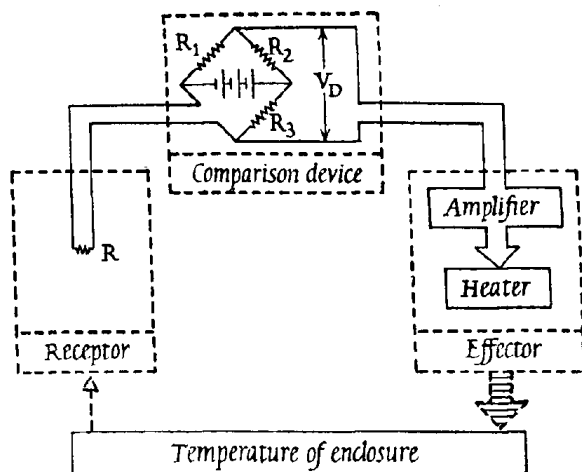


Fig. 6. Diagrammatic representation of a thermostat. The receptor consists of a platinum wire whose resistance varies with temperature. Its resistance is compared by a Wheatstone Bridge with a standard R_1 ; if R_1 is not equal to R , then a difference voltage V_D is produced which is used to actuate an amplifier and heater system; and this therefore is the effector unit. In practice, the apparatus is adjusted so that, when $V_D = 0$, the amplifier and heater produce sufficient heat to balance the normal heat losses of the enclosure at the controlled temperature.

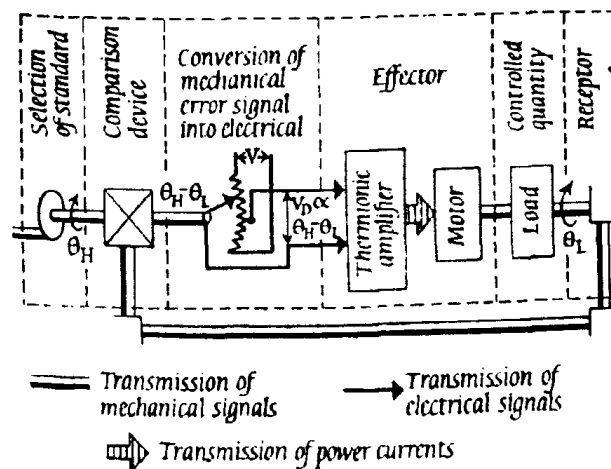


Fig. 7. General arrangement of a simplified position-control servo-mechanical system, showing the functions of the various components. In this case, the comparison standard is selected mechanically; the receptor also operates mechanically, as does the comparison device which is a differential gear system. The output from this is converted into electrical signals which, suitably amplified, operate the servo-motor. (This system is actually unstable, and will oscillate with simple harmonic motion if disturbed from the rest position; it can be stabilized by including in V_D a voltage proportional to the angular velocity of the load.)

be seen, however, that selection of standard, reception, and comparison of signals could all be accomplished electronically.

In general, therefore, we can use electronic devices to determine the operating levels of various types of machine and, if necessary, stabilize their operation at these levels by appropriate negative feedback, as shown in Figure 8. This differs from Figure 4 only in that the selection device is now more generally described, and that a finite set of actual standards is replaced by an 'information store' which can provide the equivalent of an infinite set of standards. Comparison standards can be selected

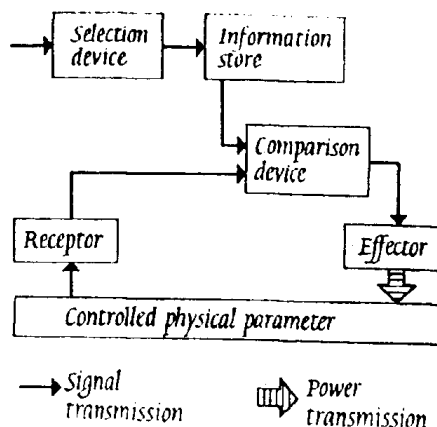


Fig. 8. Representation of a generalized electronic control system which may be stabilized by negative feedback at various determinable levels. These are represented appropriately in the information store.

by amplified electrical signals, which can thus be used to coordinate the activities of fabricating machines, materials-handling devices, and such equipment as electroplating baths. Furthermore, electrical signals from appropriate receptors can be used to indicate the extent to which various tasks have been accomplished; and thus to modify the activity of all machines in a factory in accordance both with their individual and collective performances and with general instructions for the running of the factory, provided that these are suitably coded in the form of electrical signals.

THE AUTOMATIC FACTORY

The processing of all this information, and the production of output signals to various effector devices requires the very rapid performance of large numbers of computations and comparisons. For this purpose, the modern highspeed electronic digital computing machine is admirably suited. It can process informa-

tion at extremely high speeds; it handles information in the form of electrical signals; and it can be programmed, that is to say it can be given operating instructions which enable it to deal with input information from various machines, to select appropriate comparison standards, and also to provide the output signals necessary for the operation of appropriate effector devices.

We may envisage the automatic factory, therefore, as a very large automatic machine. It will consist of assemblies of subsidiary machines (which may themselves be transfer machines and hence relatively complex) together with inspection devices, automatic materials-handling devices, and automatic assembly plant. Different factories will differ in the extent to which operating instructions are 'built-in' to the different machines, or supplied to them in the form of signals generated by other machines; and this may be done directly or, where necessary, via an electronic computing machine. This will be the nerve-centre of the factory; it will receive information from every machine in the factory; it will perform the requisite computations in accordance with its own programme; and it will transmit signals to appropriate effector devices, thus regulating the activities of machines and ensuring a smooth flow of articles through the factory. It will also be used to perform such subsidiary computations as are required by the general office staff of the factory for the organization of the input of raw materials and for decisions about the quality and quantity of output.

The programme for the computing machine will obviously be extremely complicated and therefore difficult to produce. It will, in effect, represent a coded formulation of the entire factory operations for a particular product; and it will, therefore, be different for different products. It may, however, be possible to arrange for programmes to be built up from standard sub-programmes, which may thus be combined together to give different sets of operating conditions for different products.

The Automatic Electronic Digital Computing Machine

The production of a programme for an automatic factory will, in fact, present problems similar to those encountered in

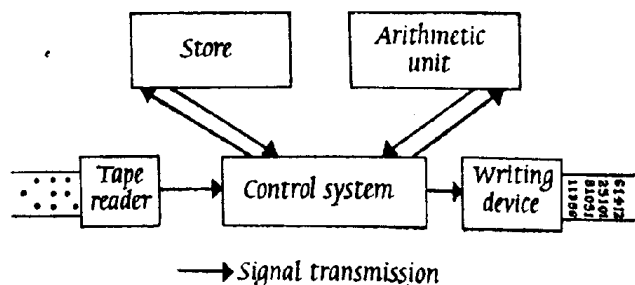


Fig. 9. Functional arrangement of a universal automatic electronic digital computing machine.

the normal programming of an automatic electronic digital computing machine, such as is represented diagrammatically in Figure 9.* It will be seen that it contains a 'reading' device which 'translates' input information in the programme into the patterns of electrical signals which represent 'machine language'. This information includes both operating instructions for the machine and the numbers and functions with which the machine will deal in the course of a computation; and it is connected directly with the control centre of the equipment. This is in effect a complicated switching system which, by suitable instructions, can be made to route information to and from the 'store' and to and from the arithmetic unit, in the same way that the complex switches in an automatic telephone exchange direct incoming information to different addresses. The operations performed by the arithmetic unit are determinable within given limits by the control unit; and this is a further example of the way in which information necessary for the assembly of a particular complex machine from given simpler machines can be stored and used in the form of electrical signals. Modern electronic computing machines operate at speeds very greatly in excess of those of tape-reading devices; so that the normal method of

*See, for example, 'Progress in Electronic Digital Computers' by R. C. Orford in *Science News* 30 (1953).

operating these machines is to transfer information from the tape to the store; the machine then operates automatically by selecting numbers from the store, in accordance with instructions contained in other parts of the store, routing these to the arithmetic unit, and suitably directing the output of this device. Finally, at the end of a computation, the control unit sends signals representing the final result to the writing device. This may be an automatic typewriter or an automatic tape-punching machine; but in any case its function is to retranslate from 'machine language' into some symbolic representation which is intelligible to human operators.

Such machines are both automatic in operation and extremely flexible. They are called universal machines because they can be so programmed as to simulate the activity of any rigid assembly of the individual units which they contain (*cf.* the discussion of universal machines given above). They are operated by electrical signals similar to those involved in the operation of fabricating and transfer machines, and they can obviously receive part of their input information directly from suitable electronic devices associated with such machines; and the output signals from them can be both recorded and used to operate suitable effector devices. As has been said before, this, together with their extremely high speeds of operation, makes these machines eminently suitable for the role of controller of the automatic factory.

Such machines are also very easily programmed to perform what are called iterative processes. This may be illustrated by considering the programming of a machine so as to enable it to divide one number by another. This is not easily done directly since these machines are most conveniently organized to perform addition, subtraction (by the addition of one number to the complement of the other) multiplication (either by multiple addition or point shift) and comparison of numbers (which may be done by subtraction). It is possible, however, to evaluate $1/a$ by calculating sufficient terms of the series defined by $Y_{n+1} = Y_n(2 - aY_n)^*$. As n is increased, the n th term of this series tends

*Provided that Y_n lies between certain values.

in the limit to $1/a$. This means that when sufficient terms have been calculated, the difference between Y_{n+1} and Y_n will become less than any chosen amount ϵ which determines the accuracy of the calculation. The calculation may thus be formulated in terms of two programmes: a 'sub-routine' which organizes the calculation of $Y_n(2 - aY_n)$; and a 'master-routine' which compares Y_{n+1} and Y_n at each stage, switching in the sub-routine if the difference is greater than ϵ , and delivering the answer when the accuracy required has been reached. The speed of operation of these machines enables them to perform such calculations with great rapidity.

Programming the Automatic Factory

This general scheme of breaking a programme down into a master programme and various sub-routines might easily be extended to the production of programmes for the automatic factory. Here the master programme would be supplied to the computing machine, and would organize the selection of various sub-routines which would determine the behaviour of individual machines or groups of machines within the factory. The sub-routines themselves might be similar to 'master programmes', merely switching on or off various machines which have their operating instructions either electronically or mechanically 'built-in'; or they might contain a great deal of information which would thus enable the computing machine to exercise detailed control of individual machine operation. In this latter case the sub-routine itself would probably be subdivided into a 'master programme' and various 'sub-routines'. There are thus three distinct levels of control within the factory. These are concerned with the organization of individual machine activity; with the coordination of the activities of machines within broad functional groups; and with the general control and coordination of all the groups of machines. The decision as to how the total amount of information required should be divided between machines, subsidiary electronic control systems, and the central computing machine will obviously depend on the type and quantity of products required, and the extent to which standard-

ization of machines and of automatic factories is possible or desirable.

In general, therefore, we may consider the automatic factory as a very large and complex machine, built up from simpler machines which, however, are themselves complex. The organization of the coordinated activity of all these machines demands an elaborate system for the transfer of information throughout the factory, together with mechanism for suitably processing this information. This is made possible by the development of electronically operated fabricating and materials-handling machines; by developments in electronics which enable any actual or required physical parameter of a system to be represented by a pattern of electrical signals; by the use of electronic computing machines which can process such patterns of signals at extremely high speeds; and by electronic and electro-mechanical amplifying systems which, actuated by the output signals of computing machinery, can control the operation and coordination of machines and groups of machines.

FUNCTIONAL THINKING

The second industrial revolution resulting from the widespread use of automatic factories will undoubtedly have effects on society and on our ways of life as profound and as far-reaching as those of the first. The development of the automatic factory with its high output of cheap consumer goods could be a great boon to mankind; on the other hand, it could distort an economy which is not as yet capable of smoothly distributing the products of the present mass-production factories; and it will obviously reduce the demand for unskilled and semi-skilled labour. The use that is made of automatic factories will depend on the forms of society within which they are developed; but however they are used, their full exploitation will pose a serious problem for the industrial designer. This will involve consideration of the optimum structural design of machines or products (which may themselves be machines) for the performance of any given function. This problem has usually been solved in the past by the more or less haphazard evolution of design, which

has generally been determined largely by the capacity and limitations of human beings. For instance, we count in the scale of ten; and this may very well be because we have ten fingers. The modern computing machine, however, using available electronic equipment, is probably most effective if it is built to work in the binary system, and most machines therefore are of this type. Similarly, automatic materials-handling devices could be built to precisely simulate human manual activity; this, however, would be most elaborate and cumbersome, and it is normally better to use combinations of hydraulically operated rams, automatic chucks, hoppers, and conveyor belts to move, say, pistons in course of manufacture from one machine to another; or to move articles by an automatic conveyor belt from station to station as is done in transfer machines. These are elementary examples, but obviously many such problems will be encountered during the next few years. The fundamental questions which will have to be asked when designing any machine or product are: what exactly is the machine or product intended to do; and what is the best way of arranging for this function to be fulfilled at the present stage of technological development. It will be readily appreciated, therefore, that this involves what might be called *functional thinking* about the problems of technology – an attitude of mind which will be immensely important during the period of development of the automatic factory.

RESEARCH REPORT

A. W. HASLETT

RADIO ASTRONOMY

A TWO-YEAR survey of 'radio stars' has lately been completed by M. Ryle, F.R.S., and his colleagues of the Cavendish Laboratory, Cambridge. The instrument used enables fainter sources to be detected and their positions to be more accurately determined than has been possible in earlier surveys. The whole sky visible from Cambridge has been covered, and in numbers alone the results are impressive. Compared with a total of about 100 sources found by all observers up to 1951, and 38 accepted as 'certain' in a list published earlier this year on behalf of the International Astronomical Union,* as many as 1,936 sources have been found. Of these, the positions of about 500 have been determined with sufficient accuracy to justify the use of large telescopes in looking in the same positions for optical objects to which the radio sources may be related. The main interest is here negative; as had been suggested by earlier surveys, the number of correspondences found is likely to remain small in proportion. Secondly, because the positions of so many sources are now known, much more can be done by studying their distribution in space statistically. This approach leads to the conclusion that the great majority of the radio sources lie outside the Galaxy, and that most are probably at greater distances than can be seen with the 200-inch telescope. A preliminary account of this work was given by Ryle in his recent Halley Lecture at Oxford. The instrument used was described by Ryle and A. Hewish at a meeting of the Royal Astronomical Society in London a week later, and the results of

*Science News 36, 116 (1955).

the survey were presented at the same meeting in the names of J. R. Shakeshaft, Ryle, J. E. Baldwin, B. Elsmore, and J. H. Thomson. The summary which follows is based mainly on the Halley Lecture.

The survey has been carried out with the new and larger radio interferometer mentioned by W. H. Marshall in *Science News* 31.* The new instrument is most conveniently approached through the previous Cambridge interferometer. This consisted of two paraboloidal reflectors and associated aerials, mounted at the two ends of an east-west line and separated by about 1,350 feet. Combining the signals from the east and west aerials causes a pattern (in sensitivity of reception) which consists of alternate maxima and minima from east to west. An extra-terrestrial source which is within the field of view of the two units individually is therefore swept through this pattern by the earth's rotation. The most intense maximum is the one at the centre – that is, in a direction at right angles to the line joining the aerials – and the passage of the source through this is timed. This measurement of time gives the Right Ascension or celestial longitude of the source. The aerials are themselves rotatable to the north or south; and Declination, or celestial latitude, is given by the position of the aerial units in which the received signals are strongest.

As in any form of radio telescope, the faintness of signals which can just be detected depends on the collecting area of the aerials. In the new Cambridge interferometer there is a change in type of unit comparable with that from the 'searchlight' type of small electric radiator, with the heating element at the centre, to the larger type which uses a cylindrical reflector. In place of paraboloidal reflectors there are reflectors in the shape of parabolic cylinders formed of stretched wires, and in place of a single aerial there is a line of aerials along the focal line of each reflector. The units are 316 feet long and 40 feet wide. They can be moved in the same way as can the cylindrical type of electric

*In his article 'Interferometry in Astronomy' which includes also an introduction to the principle of the interferometer.

radiator, except that the movement extends through the full 180° from horizontal northwards-looking to horizontal southwards-looking, so that the whole visible sky can be covered. One of these units is shown in Inset 13.

Together with the wavelength used – 3.7 metres in the present survey – the east-west spacing between units determines both the accuracy with which Right Ascension can be measured and also the largest diameter of source which can be distinguished. (The source has to be small compared with the interference pattern through which it is swept by the earth's rotation.) In the new instrument the east-west spacing has been increased to about a third of a mile. With this spacing, the directions of the brighter sources can be determined to within about 2 minutes of arc in Right Ascension, and the limit of angular diameter which can be distinguished is about 15 minutes.

There are also four units instead of two in the new instrument. They are disposed in the form of a rectangle, with the north-south spacing about 170 feet, roughly one-tenth as great as that in an east-west direction. One effect is to double the collecting area of the system, in all about 50,000 square feet. Another is that, as well as the east-west pattern of interference, there is also a north-south pattern. This is used in a more complicated way to measure the Declination, or celestial latitude, of localized sources. The accuracy of this measurement is about 12 minutes of arc. But it is also possible to operate the instrument essentially as a north-south interferometer; and, because of the smaller spacing between units, and the relatively coarse pattern of interference which this produces, sources of up to 3 degrees angular diameter can then be detected.

Of the 1,936 sources found, 30 are 'extended sources' with angular diameters between 20 minutes and 2 degrees of arc. The most intense of these lie in directions corresponding with the wide but flattened disc of the Galaxy, and one of them has been identified with a galactic nebula. However, the main interest is in the smaller sources. The positions of 500 of these have been determined within the limits of accuracy already mentioned – 2 minutes of arc in Right Ascension and 12 minutes of arc in

Declination. Those of the 75 most intense sources have been examined by R. L. Minkowski at Palomar Observatory, California, on the plates of the 48-inch Schmidt-camera survey. But only four faint galaxies have been found which, subject to confirmation with the 200-inch telescope, show peculiarities which may be significant. The positions of the remaining sources have been compared with catalogue positions. Some 80 cases have been found which suggest probable identifications, although some of these may be coincidence; and from a detailed optical study by D. W. Dewhirst at Cambridge Observatory it seems that a number of them are double or multiple galaxies which show signs of collision or interaction. But even if all are confirmed, the total is not many in proportion.

Plotted on an equal-area chart, the complete 1,906 sources of small diameter give the impression of a purely random distribution. There is no concentration in the plane of the Galaxy, as with the extended sources. This could mean one of two things: either that the sources were inside the Galaxy and at distances small compared with its thickness, or that they were outside the Galaxy. The first possibility is excluded by the results of an analysis of the number of sources with different observed intensities. This analysis has been done both for the sky as a whole and also by sectors. It leads to the conclusion that the number of sources in a given volume of space is at first constant and thereafter increases as the cube of the distance. The only assumption made is that the proportion of weak to strong sources is not itself dependent on distance. This result is difficult to explain unless it is remembered that, as distance is increased, so also is the time which has elapsed since the radio waves were emitted. But for their times of travel to be significant on a cosmic scale, the distances travelled must be very great.*

An explanation which involves time as well as distance is consistent *a priori* with evolutionary theories of the universe. It is not consistent with the 'steady state' or 'continuous creation'

*This is supported by a further line of argument which leads to a minimum distance around 500 million light-years, one quarter of the range of the 200-inch telescope.

types of theory which have lately attracted attention.* Moreover, the expectation from these theories is that the number of sources in a given volume of space should, if anything, become less rather than greater as distance is increased.

To get much further than this, it is necessary to make an assumption of some kind. A possible approach is to take as a model an intense radio source in Cygnus, which W. Baade has identified as a collision between extragalactic nebulae, and then see what sort of distance scale would be expected on the supposition that the observed faint sources are of the same absolute brightness. This is an inversion of the argument used by Ryle in his Halley Lecture, but has the advantage that the simplicity of Euclidean space, however delusive, is retained for as long as possible. From the shift towards longer wavelengths ('the red shift') in light from the source in Cygnus, its speed of recession and therefore its distance can be calculated. The latter is about 200 million light-years on the new scale of distances.* But the 'radio brightness' of the faintest radio sources is only about 1/1,000 of that of the source in Cygnus. This means that, in a Euclidean space, the distance of the weakest source would be roughly 30 times as great, that is, about 6,000 million light-years. This would unfortunately be outside the observable universe, since a nebula at that distance would be receding at a velocity greater than that of light. The answer is therefore meaningless, except to suggest qualitatively that, so far as energy is concerned, the distances of the radio sources could indeed be very great.

The argument used by Ryle starts from the opposite end. He begins with the analysis already mentioned, relating the numbers and brightness of the observed radio sources. He then asks in what way he can explain the departure of the observational curve from that corresponding with a uniform distribution in space. The reason, he supposes, is to be sought in processes

*See, for example, the article by Sir Harold Spencer Jones, F.R.S., in *Science News* 32; also pp. 22-4 of the article by W. B. Bonnor in this number.

†See, for example, *Science News* 29, 112 (1953).

which occurred at a time in the remote past when 'the observed galaxies were in close proximity and during which marked evolutionary changes are probable'. The time interval taken is the so-called 'age of the universe' – that obtained by starting from the speeds of recession and the distances of the nebulae and working backwards. This is 4,000 million years on the new scale of distances. In order to relate this time interval to the observational curve, he has to adopt some particular evolutionary model of the universe. From a 'simple model in relativistic cosmology' he calculates the absolute brightness necessary to produce the observed radio brightness. The answer is the same order of brightness as that of the source in Cygnus. In Ryle's argument, the observational curve is thus related to an assumed time interval, and the comparison with Cygnus is used only as a check. On this basis, he concludes that 'only some tens of the main class of radio star would be within reach of the 200-inch telescope', and that the distances of most of them are greater.

There is a general point which helps to explain the great distances which have been invoked and to encourage the hope that observations by radio methods will contribute usefully to cosmology. In the case of optical methods, the recession of the nebulae causes, through the red shift, a loss in energy within the limits of wavelength to which photographic plates can be made sensitive, and this loss becomes rapidly more serious as distance is increased. There is thus an observational 'cut-off' which already is being approached. But with radio methods the corresponding loss should remain unimportant at greater distances.

As well as the Halley Lecture and the two papers mentioned in the text (all in course of publication), the spatial distribution of the radio sources and its interpretation will be further discussed in a paper by Ryle, M. and Scheuer, P. A. G. (*Proceedings of the Royal Society 'A'*, in the press).

ELEMENT NO. 101

The making of element No. 101 was reported by Professor A. Ghiorso, Professor G. T. Seaborg, and others of the University

of California at a recent meeting of the American Physical Society. Its making and identification was a *tour de force* greater even than in the case of elements Nos. 99 and 100, reported a year ago.* Its chief interest is that the isotope made breaks down by spontaneous fission, with a half-life of between half an hour and several hours. The element has been named 'mendelevium' (symbol, Mv) after the Russian chemist Mendeleev 'who was the first to use the periodic system of the elements to predict the chemical properties of undiscovered elements, a principle which has been the key to the discovery of the last seven trans-uranium elements'.

Element No. 101 was made by bombarding element No. 99 with alpha-particles accelerated in the university's 60-inch cyclotron. The cyclotron was modified for the experiment to deliver an intense beam of alpha-particles (30×10^{12} a second) on to an area only $\frac{3}{32}$ inch by $\frac{1}{4}$ inch. The target atoms of element No. 99 (mass number 253), themselves with a half-life of only 20 days, were electroplated on to gold foil, on which they formed an invisible layer of less than a millionth of a microgramme weight. Atoms ejected from the target foil by the bombarding alpha-particles were caught on a second gold foil. This foil was then quickly dissolved, the gold and other impurities removed, and the remaining material separated into fractions on the basis of chemical analogies found already between the later trans-uranic elements and successive elements of the rare-earth series. These fractions consist of solutions containing element No. 99 (analogue, holmium), No. 100 (analogue, erbium), and 'a third that could only contain element No. 101 (analogue, thulium) if it were present'. The first experiments yielded nothing significant. Then, in one of the '101 fractions', bursts of energy were detected such as could be produced only by fission. This was a surprise, since breakdown with the emission of an alpha-particle had been expected. The experiments were repeated, and from a total of 17 fissions the half-life of the element was roughly estimated. Since the shortest half-life previously found for any isotope breaking down by fission was roughly 500 times

**Science News* 32, 108 (1954).

longer, the quick rate of breakdown was proof in itself that atoms of a new kind had been made. That they were atoms of element No. 101 was supported by the finding of the observed activity in the 'right' chemical fraction. The isotope which it was expected would be made was mendelevium-255, which it was predicted would break down by alpha-particle emission with a half-life of about 5 minutes. The theory of this type of breakdown has proved remarkably reliable in other cases, and it is suggested that a different isotope - probably mendelevium-256 - has been produced. Starting from uranium-238, as many as 9 units of charge and 18 of mass have now been added by stages in making this newest of the elements.

Ghiorso, A., Harvey, B. G., Choppin, G. R., Thompson, S. G., and Seaborg, G. T. *Physical Review* (in the press).

ROYAL SOCIETY CONVERSAZIONE

The phenomenon of 'industrial melanism' in moths, first noticed in Manchester about a century ago, was illustrated by Dr H. B. D. Kettlewell of the Genetics Laboratory, the University Museum, Oxford, at a *Conversazione* held in May by the Royal Society in London. In industrial areas, and in districts to which industrial smoke is carried by the prevailing wind, some 50 species of the larger moths are in course of changing their initial light wing pattern to all-black colouring. Within 50 years, the population of the 'Peppered Moth' has changed in some areas from being 99 per cent light to being 99 per cent dark. Research on the problem is being done in several directions. The proportions of light and dark forms have been correlated with degree of smoke pollution. From the release of numbers of marked black and light moths in the industrial Midlands, and their subsequent recovery, it has been confirmed that where tree trunks are black there is an advantage for survival in black colouring; the proportion of black moths recovered has been about twice as great. There is also a biochemical aspect; the larvae of a black form appear, in some cases, to be better adapted to feeding off soot-covered leaves.

The uses of high-speed cinematography in physiological research were illustrated in a film made at 1,000 to 2,000 frames per second by John Hadland of Chipperfield, Herts, and shown by Dr D. A. McDonald of the Department of Physiology, St Bartholomew's Hospital Medical College. Dr McDonald has used the method extensively to measure the velocity of blood flow in the arteries of experimental animals. He has found, for example, that the flow in a dog's femoral artery accelerates from zero to about 1 metre per second in about one-twentieth of a second, and that there is always then a period of backflow, during which blood flows towards the heart. Other sequences in the film showed fast movements which depend on nervous reflexes, including the righting of a cat when dropped upside down. Stills from a similar high-speed record are reproduced as Inset 14-21. They show that the classical explanation in terms of an initial slight twist and adjustment of the moment of inertia by movements of the legs and tail is an over-simplification. An essential element in all records is that the head and shoulders are turned first, and that the turning of the rest of the animal follows. There is a contrarotation of the tail, as can be seen from the pictures. An analysis of the movement is to be published.

MAN-CARRYING CENTRIFUGE

The limited ability of the human body to withstand accelerations and decelerations many times that of gravity is a well-known problem in the design and manoeuvre of high-speed aircraft. To assist in its study, a man-carrying centrifuge which provides accelerations up to 30 times that of gravity has been installed at the R.A.F. Institute of Aviation Medicine, Farnborough, and is illustrated in Inset 22-6. It consists essentially of a horizontal arms (Inset 22) rotating under a controlled cycle of acceleration and deceleration, which can be repeated within ± 2 per cent over the whole cycle. Alternatively, manual control can be used. Figure 1 shows the general arrangement. The arm, about 60 feet long, is centrally mounted (Inset 23). At either end, a car is suspended in which the subjects of the experiment sit

and are observed (Inset 25). The cars can be attached at distances of 10, 15, 20, 25, or 30 feet from the centre; in the outermost position the greatest speed reached by the car is 115 miles per hour. Two observers, one for each car, are seated back to back near the centre of rotation (Inset 24). Information from the two halves of the equipment is passed continuously to recording stations in a laboratory. There are teletalk circuits which include both subject and observer. Both the external 'controller' – who has a full view of the rotating arm – and the 'recorders' in the laboratory can also enter the circuit. More important for experiment, electro-medical recordings, which may be of low amplitude, are passed from the whirling cars to the laboratory. A separate slipping column (Inset 26) is used for the teletalk and recording circuits. It is fitted with 60 silver slippings, divided into three banks of 20 each, and 30 mercury troughs divided into five banks of 6 each. As well as its main use for work at high values of acceleration, the equipment can be used also to study ability to perform simple tasks at comparatively low accelerations. For this purpose, the cars are removed and replaced by box-like structures at the ends of the arms.

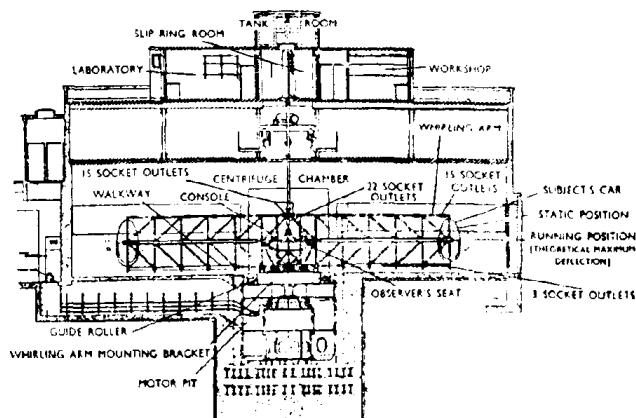


Fig. 1. General arrangement of man-carrying centrifuge.

GREENLAND ICE

Readers of the account by E. F. Roots of the Norwegian-British-Swedish Antarctic Expedition 1949–52, published in *Science News* 26, will recall that, beneath the mounting surface of ice, the underlying bedrock was shown by seismic sounding to descend below sea level at places more than 400 kilometres inland from the expedition's base at Maudheim. Preliminary results of geophysical work by the British North Greenland Expedition show that across much of the eastern part of northern Greenland no definite reflections at all could be detected from an ice-bedrock surface. There was a dividing line, running roughly north and south at about 44° W., east of which reflections ceased to be obtained until the extreme eastern edge of the ice-sheet. At the last stations yielding reflections, the level of the bedrock was close to or below sea level, and to the west of the line the bedrock level rose. No conclusions are drawn, but the implication that much of Greenland is 'not land' seems reasonably clear.

Bruce, R. J. M., and Bull. C. *Nature*, 175, 892 (1955).

CORRESPONDENCE

GAMES WITHOUT A STOP

In *Science News 16* it was pointed out that in draughts or chess some rule was necessary to stop the game after a finite number of moves. Otherwise the players might 'dodge about' indefinitely.

But one naturally thinks; in a position where white can win, he is foolish merely to dodge about. In rational play, dodging about could only occur where neither side saw a way of forcing a win. Thus, even without rules to stop dodging, it should be possible to analyse all positions into wins, draws, and losses for white.

To take chess as an example, we could give a procedure as follows. There are only a finite number of *positions*. We can make a list of all the positions in which white has just moved and delivered checkmate. Attach the label $A(0)$ to all such positions. Let $B(1)$ label any position in which white is to move and can give mate in one move. We can list all positions $B(1)$ by going through all the $A(0)$ positions, and listing all the ways in which white's last move can have been made. Continue to work back. If black is playing rationally, he will not move into a $B(1)$ position unless he has no better choice. Accordingly we can make a list of positions $A(1)$ in which whatever black does, he is bound to leave a $B(1)$ position. We then make a list of positions $B(2)$ in which white has a move that will produce an $A(1)$; these are positions in which white can mate in two moves, but not in one move; we do not include in $B(2)$ any positions already listed in $B(1)$.

Continuing in this way, we reach a list $B(n)$ of all the situations in which white can mate in n moves but not in less, and the positions $A(n)$ in which black has to move and has no option but to leave a $B(n)$ or something worse for himself. Thus an $A(n)$ is a position in which black cannot escape defeat, but can prevent white from mating in less than n moves.

Since each stage of this process adds new positions to the

tree, and since there are only a finite number of positions, the process must eventually terminate. Any game which is played rationally, and leads to a win for white, must appear in this tree.

In the same way, we can construct a list of all positions in which rational play leads to a win for black. In the remaining positions, rational play can lead only to a draw. Thus we have a finite procedure for attaching win, draw, or loss to every position.

W. W. SAWYER

Mathematics Department,
Canterbury College,
Christchurch, New Zealand

D. W. Davies writes:

Most treatments of games use the 'strategy' approach, in which the strategies can be chosen in advance by the two players; then the outcome can be deduced from the rules of the game. In this approach, plays which cycle indefinitely must be eliminated.

The treatment I gave in *Science News 16*, which only applies to games of perfect information, can perhaps be extended to some cases where closed loops of moves are allowed. By applying the maximum and minimum rules, inequalities are in general obtained for the values of positions in a closed loop. In some cases the values are given precisely by these inequalities.

In the case of only three possible outcomes, $+1, 0, -1$ (or win, lose, or draw), if the value is not given precisely, then 0 (or draw) must be a possible value. I believe that Sawyer's suggestion is to make the value 0 in such cases.

In the case of more than three possible outcomes, it is not so easy to give a value to these positions. Special rules would be required for any given game.

In either case, I regard the assigning of a value to these positions as a special feature of the game, analogous to a stop rule, not as part of the theory of games. But, sometimes, the theory of games may give an exact value — a point I did not mention in my article.

SOME BOOKS RECEIVED

FRONTIERS OF ASTRONOMY by Fred Hoyle (London: William Heinemann Ltd, 1955), pp. 360, 25s.

Mr Hoyle's fertility of ideas is equalled only by the conviction with which he expounds them. In his present book, more substantial than its predecessor, he works outwards from earth to receding nebulae, and on the way disposes of the origins of ice ages, planets, stars, chemical elements, life, and the nebulae themselves. Space is taken for some background physics and the presentation of argument, but reservations and alternatives are on the whole sparsely represented. There is also much recent material, not easily accessible. The result is a mainly consistent picture covering a wide area of knowledge, theory, and speculation. It is a remarkable performance; but how many readers, one wonders, will be able to distinguish between one component and another?

THE ORIGIN OF THE EARTH by W. M. Smart (Harmondsworth: Penguin Books Ltd, 1955), pp. 224, 2s 6d.

Professor Smart's book, first published in 1951, offers a simple but critical discussion of a single problem in astronomy. His summing up is that, while both a common origin of the solar system and current estimates of the age of the earth can be accepted with some assurance, 'all the theories proposed up to date as to the mechanism by which the solar system has come into being fail to carry conviction'.

BIRD NAVIGATION by G. V. T. Matthews (Cambridge University Press, Monographs in Experimental Biology, 1955), pp. 141, 12s 6d.

Dr Matthews wrote in *Science News* 23 on 'Problems of Bird Orientation'. In his present monograph he gives a fuller account than was then possible and includes also the results of important experiments which strengthen greatly the evidence for some form of navigation based on the sun. For this there is now evidence from pigeons, two species of gulls, and the Manx Shearwater, that for the latter species being close to conclusive. Dr Matthews's further suggestion

that navigation is based on memory of sun arcs, plus an internal chronometer, remains an hypothesis which, if difficult either to accept or to test adequately, appears at least to be less unlikely than any other which has yet been put forward.

THE FOSSIL EVIDENCE FOR HUMAN EVOLUTION by W. E. Le Gros Clark (Chicago University Press; U.K. agents, Cambridge University Press, 1955), pp. 181, 45s.

A review of human origins, and in less detail those of the anthropoids, is preceded by a critical introduction on terminology and methods, one of the author's points being that 'man' is a dangerous word. Points which emerge from the body of the book are the tendency of different parts of the skeleton to show different degrees of specialization, and the undesirability of multiplying generic names on evidence which usually is scanty. Lucid in substance, Professor Le Gros Clark's review could have been more easily written.

ANALYTICAL CHEMISTRY: THE WORKING TOOLS edited by C. R. N. Strouts, J. H. Gilfillan, and H. N. Wilson (Oxford University Press, Geoffrey Cumberlege, 1955), 2 vols., pp. 1067, £5 5s.

This unique publication embodies the cumulative experience of Imperial Chemical Industries Ltd since 1928, when an analytical chemists' committee was set up within the company to standardize analytical methods so far as possible. With the minor exception of a few techniques mentioned for potentialities, it is claimed for it that 'all the methods . . . described have been found accurate and reliable in the hands of a number of different analysts working in different laboratories'. Directions and explanations are given with meticulous clarity. For both reasons, it should be an ideal reference book for any laboratory which is called upon to undertake analyses of any wide range of materials.

ATOMIC ENERGY RESEARCH AT HARWELL by K. E. B. Jay (London: Butterworth's Scientific Publications, 1955), pp. 144, 5s.

The 1952 report on the work of the Atomic Energy Research Establishment, Harwell, is carried forward in two ways. Part I describes advances in the main programmes, and is meant for general

readers (well informed ones in practice). Through no fault of the author, this will look dated after Geneva. Part II describes selected researches, mostly fundamental, in the divisions of physics, chemistry, chemical engineering, and metallurgy, and is meant for scientists reading away from their own work. This will wear better, and is impressive as suggesting the large number of directions in which work at Harwell is likely to be useful. Examples range from cosmic ray physics to corrosion and the distribution of smoke from tall chimneys.

OXFORD JUNIOR ENCYCLOPAEDIA: VOL. VIII, ENGINEERING. (Oxford University Press, Geoffrey Cumberlege, 1955), pp. 496, 30s.

On the supposition that users of encyclopaedias are unlikely to be specialists in subjects looked up, there should be an advantage in one which is authoritative, directed primarily at school libraries, and simple but with avoidance of writing down. With the reservation that key (but difficult) points appear in some cases to be dealt with by omission, it may be guessed that more people than would care to admit would find this a better place to begin than most.

ABOUT OUR CONTRIBUTORS

W. B. BONNOR is Lecturer in Applied Mathematics in Liverpool University. He took a degree in mathematics at Birkbeck College, London, after having started his scientific career as a chemist. At present he is interested in general relativity and cosmology, and he has worked also on unified field theory and mathematical chemistry.

J. B. COLLINS was trained as an illuminating engineer in the research laboratories of the General Electric Company. In 1939 he joined the Admiralty to study problems of searchlight performance and visibility. At the end of the war he helped to establish a laboratory for photographic and instrument research for the Royal Naval Scientific Service, and was one of a three-man team responsible for developing underwater cinematography in this country. He is now concerned with research on the more fundamental problems of the lighting of buildings at the Building Research Station at Watford. He is a Fellow of the Illuminating Engineering Society.

J. B. CRAGG has been Professor of Zoology, Durham Colleges in the University of Durham, since 1950, having previously been Reader in charge of the department from its foundation in 1946. Aged 45, he worked formerly on comparative physiology in the University of Manchester, and in the Department of Zoology, University College of North Wales. During the war he was seconded to the Ministry of Agriculture and later became a member of the Agricultural Research Council's Unit of Insect Physiology.

KATHLEEN LONSDALE (*née* Yardley), F.R.S., is Professor of Chemistry and head of the department of Crystallography at University College London. She was born in S. Ireland, the youngest of ten children, and educated in England, first at State Schools and then at Bedford College (London University), taking her degree in physics in 1922 at the age of 19. She joined the research team directed by the late Sir William Bragg at University College London, and then at the Royal Institution, and remained there until 1946, apart from short absences when she married in 1927, and when she had three children. She was one of the first women to be elected to the Fellowship of the Royal Society.

J. D. TURNER is a chemistry graduate of Sheffield University. During the war, as a technical officer in R.E.M.E., he was engaged in work

on radar and electronic equipment. After some time spent as a senior executive with Unilever, he returned to Sheffield University, where, as Staff Lecturer in Science, he is responsible for all scientific and technological courses run by the Extramural Department. Particular interests, the history of science and technology.



Am 1, 246
N55-39.