

SCIENCE NEWS

EDITED BY ARCHIE AND NAN CLOW

46

PENGUIN BOOKS

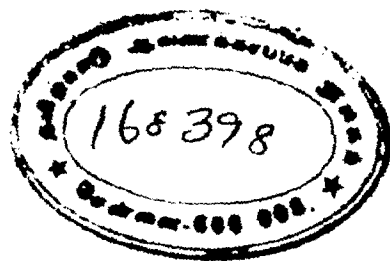
Penguin Books Ltd, Harmondsworth, Middlesex
U.S.A.: Penguin Books Inc., 3100 Clipper Mill Road, Baltimore 11, Md
AUSTRALIA: Penguin Books Ltd, 262 Whitehorse Road,
Mitcham, Victoria

First Published November 1957

DL

Am 1, ~ 46

NS 2-46



Made and printed in Great Britain
by The Campfield Press,
St Albans

CONTENTS

THE TWILIGHT OF PARITY CONSERVATION J. R. Singer	7
NEW APPROACHES TO THE STUDY OF BACTERIA D. M. Adams	19
MINERAL DRESSING E. J. Pryor	35
THE ORBIT OF A SATELLITE W. C. Vaughan	43
NUCLEAR COLLISIONS PART II P. E. Hodgson	58
AVIATION MEDICINE G. W. G. Sharp	74
RESINS IN EMULSION H. Warson	88
RESEARCH REPORT A. W. Haslett	107
CORRESPONDENCE	119
SOME BOOKS RECEIVED	121
ABOUT OUR CONTRIBUTORS	126
RECENT ISSUES	128

19 illustrations are between pages 64 and 65

THE TWILIGHT OF PARITY CONSERVATION

J. R. SINGER

EARLY this year Columbia University announced that experiments performed in co-operation with the U.S. Bureau of Standards low-temperature laboratories left little doubt that parity was not conserved – at least for three interesting nuclear reactions. What this may mean for physicists is an abandonment of many years of theoretical description, and a reconsideration of much of their work on nuclear theory. It is not easy to convince scientists that a well-accepted principle is wrong, and in the following article I shall try to describe the evidence that has overwhelmed some of the older conceptions.

DEVELOPMENT OF QUANTUM THEORY

In 1949 Enrico Fermi said: 'The concept of parity arises from the fact that neither the Hamiltonian (the mathematical operation giving the energy of a system of particles) nor the energy of an isolated system of particles (e.g., an atom or nucleus of an atom) will change if we invert simultaneously through the origin the co-ordinates of all the particles of the system.'

To Fermi and all physicists it was clear that space is isotropic, and it made no difference whether we observed an atom from the left side or the right side. Hence a change in position should not alter the net energy.

Fermi, continuing, said: 'Parity is a quantum-mechanical concept, we do not know of any classical analogue.'

We must, therefore, adopt a modern or wave mechanical description of particles rather than Newton's physics if a description of parity is to be realistic. It is difficult to condense so involved a subject, but an attempt follows.

Quantum mechanics owes its origin mainly to Max Planck, who in 1900 was forced to the conclusion that light is emitted and absorbed in packets of energy called quanta. This was a startling concept to contemporary physicists, yet they could find no way to circumvent Planck's quantum explanation of the way that energy was distributed in the light waves given off by heated bodies. A few years later Albert Einstein showed that the photo-

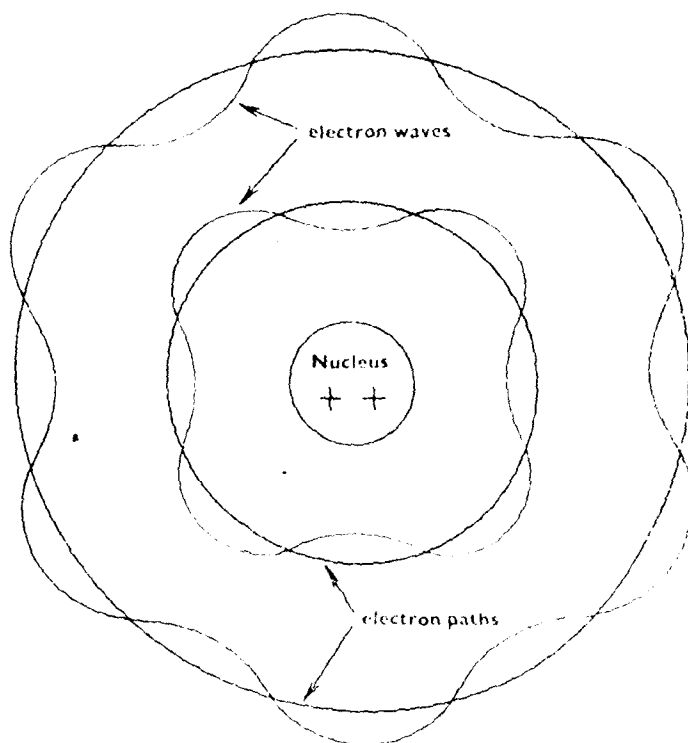


Fig. 1. Bohr model of the atom in which two electrons are represented as forming energy conserving standing waves as they revolve about the nucleus.

electric effect was easily understood if light energy was considered to be carried by particles rather than waves. These particles are called photons or quanta of energy. At this point, physicists began to believe in a discontinuous energy distribution, and developments occurred rapidly.

In 1913 Niels Bohr devised an ingenious system of positive charges surrounded by a planetary system of electrons to describe atomic structure. Each of the electrons was permitted to take a path around the nucleus comprising an integral number of wavelengths. Since the electron would continue to establish the same wave pattern, no energy would be radiated by the 'standing wave'. A representation of this atomic model is given in Figure 1. The implications of such a structure were advanced further by Louis de Broglie who, in 1924, suggested that an electron acts either entirely as a wave or entirely as a particle. So radical a thought might not have rapidly had wide acceptance, but at this time Davisson and Germer were experimenting with electrons penetrating matter, and found that the electrons provided the same sort of interference patterns as do X-rays. Thus de Broglie's speculation was accepted almost immediately. But if electrons could act as waves, it should be possible to describe their paths by a wave equation, and in 1926 Erwin Schrödinger derived such an equation. This description of electron motion replaced Bohr's fixed paths for electrons encircling the atomic nucleus with a rather hazy set of conditions associated with the Schrödinger equation. (See Figure 2 for this atomic model.) The degree of 'haziness', or uncertainty in the electron paths was explained by Werner Heisenberg, in 1927, by his uncertainty principle. This principle states that an exact measurement of the energy of a particle and an exact knowledge of the simultaneous time are incompatible. Furthermore, the more precisely the time is known, the less precisely the energy at that time may be determined. In addition to the relationship between time and energy, position and momentum also act as interacting measurements; both being non-determinate simultaneously. In practice, this indeterminacy or uncertainty is an extremely small amount; a micrometer measurement will in no

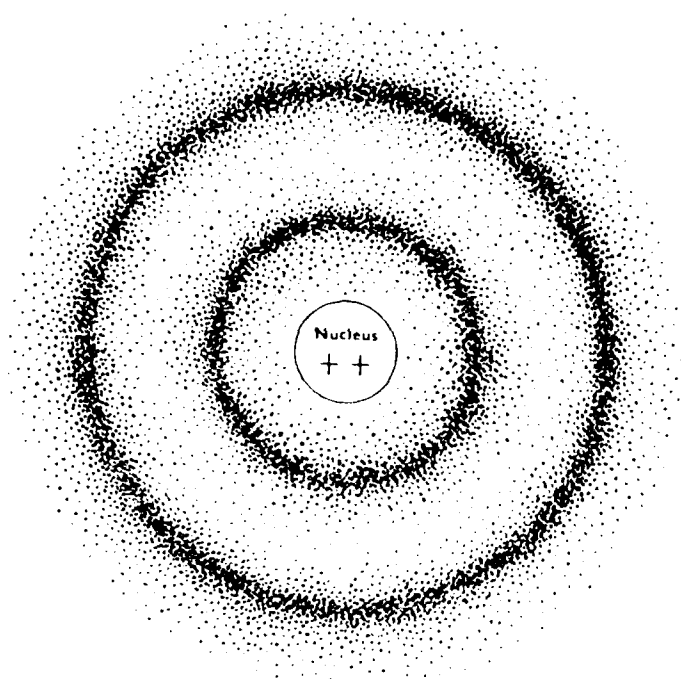


Fig. 2. Wave mechanical model of an atom showing the most probable positions for two electrons. The denser regions represent positions in which it is more likely that electrons will be found.

case be affected. None the less, to date, every successful atomic theory has incorporated the uncertainty principle.

One aspect of quantum theory that is not needed to explain the grosser aspects of matter, but to complete the theory for fine details, is the theory of spin. In order to understand narrowly spread energies of light waves emitted from heated matter, Samuel Goudsmit, now head of Brookhaven National Laboratories, Long Island, N.Y., and G. E. Uhlenbeck, now professor of physics at the University of Michigan, postulated in 1925 that electrons possessed 'spin'. Exactly what spin means in the

macroscopic world no one knows, but the effect of electron spin is to provide a magnetic electron. The particles in the nucleus are also magnetic; hence forces exist which depend upon whether, relative to the rest of the atom, the electron is 'up' or 'down'. The difference in energy between the up or down orientation gives rise to a difference in the energy of light waves emitted by the atom. Thus Uhlenbeck and Goudsmit justified their spin hypothesis. Later, it was discovered that all particles possess spins in integer multiples of the spin of an electron.

We have now reviewed the highlights of the quantum theory. The real test of a scientific theory is agreement with experiment — even better, a theory is most successful if the outcome of an unperformed experiment can be predicted successfully. Quantum theory has done just that many times, because, in addition to explaining the energy spectrum of light emitted from a heated set of atoms, it has been used to predict the energy of particles emitted from radioactive materials such as radium, plutonium, and thorium.

In spite of the demonstrated success of the quantum theory in the atomic field, it must also describe how billiard balls, stones, and other large particles act with energies and constraint forces imposed. This property, a description of the classical or Newtonian particle as well as the atomic or sub-atomic particle, is known as the correspondence principle. In both size domains, atomic and polyatomic, certain measurable characteristics of the system composed of the particle and the forces are invariant. For example, both billiard balls and atoms satisfy the conservation of energy principle. For billiard balls, energy is given to the table as the heat of friction, and to the air as sound vibrations, but no energy is ever completely unaccounted for, though it may no longer be useful. Where atoms are involved, energy of motion is often given off in light particles, or even in a change in mass, following Einstein's relation stating that energy is proportional to mass, but, even if invisible, all energy is conserved — almost undetectable neutrinos carry off energy. Likewise, momentum is conserved in both the macroscopic everyday world and the tiny atomic world described by mathematics.

CONSERVATION LAWS

In a nuclear reaction, certain conservation laws appear that do not apply to large classical particles. This may be shown by examining what happens when a photon strikes an isotope of lithium, containing three protons and four neutrons in the nucleus, and hence written Li^7 . If the photon, which is a quantum of light energy, has enough energy, an electron is driven out of the nucleus, and since the charge increases by one positive unit, the lithium becomes beryllium as illustrated in Figure 3. This is

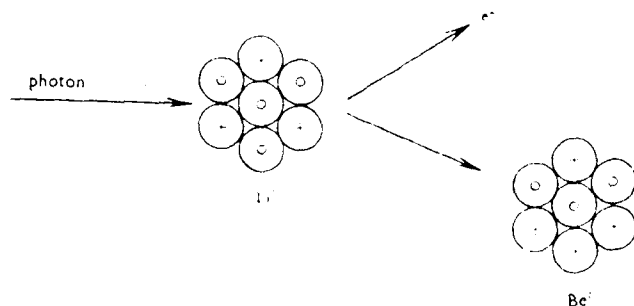


Fig. 3. Representation of the transmutation of lithium into beryllium. Charge, energy, momentum, and spin are conserved. Parity conservation is not tested here, but there is no violation unless the reaction has a preferred orientation, which cannot easily be ascertained.

an achievement worthy of a great alchemist. An element is transmuted. The interesting point here, however, is not the change, but the conservation laws – what has remained invariant? By careful measurements, it is found that energy and momentum are conserved. Electric charge is conserved – this is what changed the element. Parity is conserved because the directions of the incoming photon and the outgoing electron may, as far as we know, be reversed. However, parity conservation appears as a result of the equivalence of any direction in space

in the equations, rather than by the exhibition of any experimental verification from this one experiment. One other conservation law might be observed; this is spin conservation. Particles have either even spin or odd spin. The sum of the spins of the lithium and the photon is equal to the spins of the electron and the beryllium.

BEHAVIOUR OF K PARTICLES

Many thousands of nuclear reactions can be demonstrated which satisfy the five conservation laws, energy, momentum, parity, charge, and spin; but in 1956 the attention of nuclear physicists was directed to a recently discovered K particle, which exhibited unusual behaviour. This particle, which is called a meson (because its weight is less than a proton and more than an electron), breaks down in the two ways shown in Figure 4. Sometimes two smaller mesons, called π mesons, are

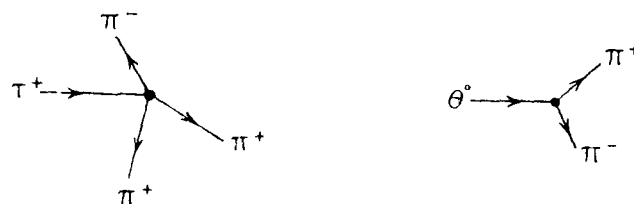


Fig. 4. Two ways in which K particles decay. One particle (K^+) decays into three π mesons; the other (K^0) decays into two π mesons. These were originally observed by means of bubble chambers and photographic emulsions combined with the output of nuclear accelerators.

the decay products, and sometimes three π s appear. Such behaviour does not conserve parity, which demands that the decay process be unique.

In view of the fact that parity conservation is satisfied by many other decay schemes, the first proposal advanced to explain two and three particle decay of K mesons was that two types of K particles existed. If that were so, then parity would

be conserved, and all the existing theoretical laws would then be correct. However, no one was able to detect the slightest difference between a K particle decaying into two π mesons and a K particle decaying into three π mesons.

At that point, Professor Tsung Dao Lee, a professor of theoretical physics at Columbia University, New York, and Professor Chen Ning Yang of The Institute for Advanced Study at Princeton, advanced another hypothesis — that parity is not conserved, and that perhaps sufficiently sensitive experiments would show this lack of reflexion invariance. They then proceeded to outline two experiments that would constitute a conclusive test.

For one test, it was necessary that radioactive cobalt be permitted to decay while a powerful magnet aligned the magnetic cobalt atoms in one direction. The electrons given off by the radioactive cobalt were to be counted, and if more particles came off in one direction than another, this would violate the conservation of parity law. Since this experiment had to be done at a temperature of close to -273°C , Drs Ralph P. Hudson, Ernest Ambler, Raymond Hayward, and Mr Dale Hoppes, all scientists in the U.S. Bureau of Standards laboratories, co-operated in setting up the experiment. They were also aided by Professor C. S. Wu of Columbia University.

The experimental arrangement is shown in Figure 5. The number of electrons per second given off by the radioactively decaying cobalt is always constant. Scintillation counters* surrounded the cobalt so that the relative number of electrons in a fixed direction could be counted. When the magnetic cobalt was aligned in one direction by the powerful electromagnet, the number of electrons emitted in the fixed direction was recorded. When the magnet was turned off, another value for the number of electrons emitted in that direction was found, and finally, when the magnet was reversed, a third value of electron emission was noted. There could be no conclusion except that the direction of electron emission depends upon the magnetic alignment of the cobalt atoms since electrons appeared more fre-

* E. S. Williams, 'The Scintillation Counter', *Science News* 36.

The Twilight of Parity Conservation

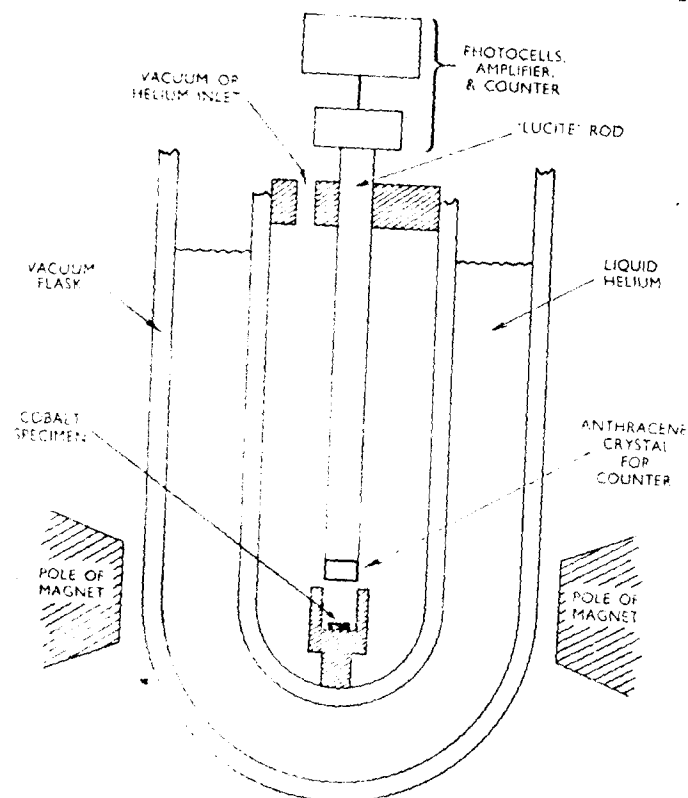


Fig. 5. Diagram of the experimental arrangement used to count the electrons given off by a magnetically oriented radioactive cobalt source. The plastic rod carries the light given off by the scintillation crystal to photocells and counters.

quently along that direction than in any other. That is, the electrons preferentially have a path along the spin direction of the decaying cobalt atoms. Before the results of this experiment were reported, another group of experimenters completed a different parity conservation test.

This second experiment also required observation of asymmetry of the decay product. In this case, Professor Leon Lederman, Dr Richard Garwin, and Mr Marcel Weinrich of Columbia University looked at the electrons given off by the π meson that disintegrates into a μ meson and a neutrino. The orientation of the meson spin relative to its motion was determined by observing the electrons given off when the μ mesons decayed. The experimental arrangement is shown in Figure 6. A cyclotron generates π mesons which decay into μ mesons. The μ mesons collide with a carbon block where they come to rest. A

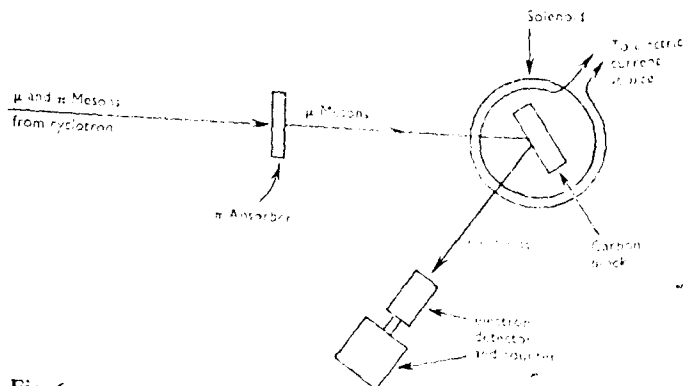


Fig. 6. μ and π mesons from a cyclotron are separated by absorbing the π s. The μ s end up in a carbon block where they are aligned magnetically. The electrons resulting from their decay are counted as a function of the magnetic alignment which is proportional to the magnetic field (or current) in the solenoid.

magnetic field aligns them since they have a magnetic moment. The number of electrons emitted in a fixed direction as they decay was found to depend on the alignments. Here, once more, the direction of electron emission depends upon the direction of spin alignment.

Both the experiments show that parity is not conserved, because in each case a mirror image of the observation apparatus

would result in a different number of observed particles. It is perhaps more startling if a spinning particle is considered as a large ball. Suppose, now, that this particle decays, emitting electrons mostly in the upward direction. A mirror image of the particle shows it rotating in the opposite direction, yet it still emits electrons mostly upwards in the image as illustrated in Figure 7. This implies that a mirror image observation is only equivalent to a non-imaged situation when the spin or left to right situation is not related to the top or bottom events. Since three known experiments demonstrate that this is not so, the principle of parity conservation does not satisfy all quantum mechanical situations.

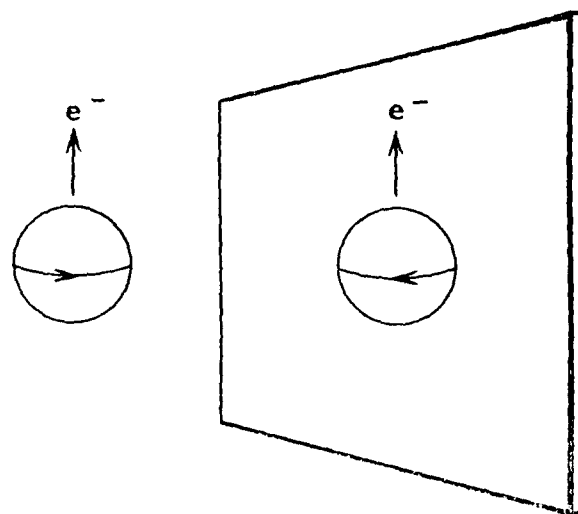


Fig. 7. A mirror image of a spinning particle emitting an electron in a preferred direction. Parity is not conserved because the image view is not identical with the direct view. Although a spinning magnetic particle would seem to provide the same lack of mirror invariance, it does not, because the north and south poles of a magnet are arbitrary names, whereas particle emission determines an absolute direction.

Whether the work on parity laws leads to a more complete physical theory or not, it is a fascinating aspect of the teamwork achievements of modern physics. One of the most interesting consequences now currently proposed by Professor Yang is that the neutrino and anti-neutrino differ by travelling as right-handed screws and left-handed screws. That is, the spin of the neutrino is clockwise when facing in the direction of its motion, while the spin of the anti-neutrino is counter-clockwise. This is, then, perhaps the only way in which particles without charge can be distinguished from their anti-particles.

Since the above was written, J. Friedman and V. L. Telegdi of the University of Chicago have found that π mesons decay to μ mesons with the axis of spin along the direction of motion. This is another example of the non-conservation of parity.

The author would like to thank Drs T. D. Lee, C. N. Yang, R. Hudson, M. Lederman, and M. Weinrich for supplying unpublished data.

FURTHER READING

J. Friedman and V. L. Telegdi: 'Nuclear Emulsion Evidence for Parity Non-conservation in the decay chain $\pi^+ \rightarrow \mu^+ + e^+$ ', *Physical Review*, 1957, **105**, 1681.

T. D. Lee and C. N. Yang: 'Parity Non-conservation and a Two Component Theory of the Neutrino', *Physical Review*, 1957, **105**, 1671.

NEW APPROACHES TO THE STUDY OF BACTERIA

D. M. ADAMS

BACTERIOLOGY is undergoing a transition. A vast literature exists describing what bacteria do, but there is little that describes what they are, and how they function as autonomous living beings. During recent years, however, an increasing interest has been taken by biochemists, chemists, and physicists in the bacterial cell, and a fairly clear outline of its structure is now emerging. It is the purpose of this article to describe briefly present views on the structure of the bacterial cell, and to give some idea of the methods being used in its elucidation.

STRUCTURAL FEATURES

Bacteria consist of a cytoplasmic membrane, which is semi-permeable and contains the cytoplasm, nuclear apparatus, and other cellular inclusions. These together form a unit, the protoplast, which can be separated from the cell wall and is capable of an autonomous existence. Surrounding this membrane, and probably secreted by it, is the cell wall, a structure of considerable rigidity which is responsible for maintaining the shape and mechanical integrity of the cell under osmotic pressures that may be as high as 20 atmospheres. Until recently the chemical nature of the cell walls was unknown, although it had been possible to make speculative guesses on the basis of, for example, immunological reactions. There remains one more part of the cell to describe: outside the cell wall proper there is often a slime layer, or capsule, the thickness of which can usually be varied by altering the composition of the growth medium. This part of the cell, being the outermost, has important immunological reactions, and it may play other roles as well. Like the rest of the

cell wall, however, the capsule does not appear to be an essential part of the cell. Indeed the cell can continue to exist without it.

To study the various parts of the cell we must either study it *in situ* by a method we know to be specific for a certain type of property, or else we must take the cell apart and study the pieces. Both of these approaches have paid dividends. The cell walls may be separated from the protoplasts by mechanical disintegration. This is achieved by shaking a concentrated aqueous suspension of cells with very small glass beads, followed by washing and differential centrifugation until the preparation is pure. Alternatively, ultrasonic waves may be used to bring about disintegration. It is customary to check the purity of the final preparation under the electron microscope. So examined the cell walls are seen as collapsed balloon-like structures, which can easily be distinguished from the more absorbant electron-dense cytoplasmic material.

To prepare the delicate protoplasts a much more refined technique must be used. First the cells must be immersed in a medium of sufficiently high osmotic pressure to enable the protoplasts to remain intact when free from the cell walls. The cell wall is then removed by treatment with the enzyme, lysozyme, thus leaving free protoplasts, which, if suitably protected osmotically and mechanically, continue to grow and divide. If, however, the osmotic pressure of the surrounding medium is reduced the protoplasts lyse and the membrane is left as an empty case which, when itself disintegrated, forms particles that in the case of Gram positive species are of lipoprotein and form something like 10 per cent of the cell weight.

METHODS OF INVESTIGATION

Having described in outline the structural features of the bacterial cell and methods of splitting it up for study, we must now look at some of the more important methods that are in use and see the type of results they can be expected to give, and some of those they have already given. Working from the outside of the cell inwards, we begin with the slime layer or capsule.

Electrophoretic studies

When a charged particle of colloidal dimensions is suspended in a liquid between two charged electrodes, it moves towards the electrode with a charge opposite to its own. Thus, bacteria with an overall negative charge at physiological pH's, owing to the ionization of carboxyl and phosphate groups on their surfaces, move towards the positive electrode with a measurable mobility. This is known as electrophoresis and it is especially valuable in the study of bacteria because we know that by taking these measurements we are observing a property of the outermost layer of the cell, often a slime layer, which is the first layer met by anything approaching the cell for absorption or ingestion. Briefly the problem is this: given the mobility of the cell - which is related to the surface charge - and its variation under certain chosen conditions, how to determine the surface composition.

One approach that is often useful is to study the variation of mobility with the pH of the suspending medium. This involves, in effect, a titration of the surface groups. Figure 1 shows the result of such experiments upon two different species of bacteria. As the pH decreases and the suspension becomes more acid, negative groups upon the surface of the cells are neutralized by the hydrogen ion of the acid being added and the charge gradually reduced. If the charge can be reduced to zero and then reversed it is often an indication that positive groups are present on the surface, which hence may contain some protein or other constituent having basic groups. The results of this kind of work on cells may be compared with similar work on non-living surfaces of biological substances such as oil droplets covered with lipid or polysaccharide. But, although quite a number of these so-called model surfaces are available for comparison, this would still not constitute an infallible method, because there may be some overlap between the behaviour of the various classes. A more specific method has been developed recently, however, which is based upon a fundamental property of the ionic groups of the surface, namely, their polarizability. This property is a measure of the ease with which the positive and

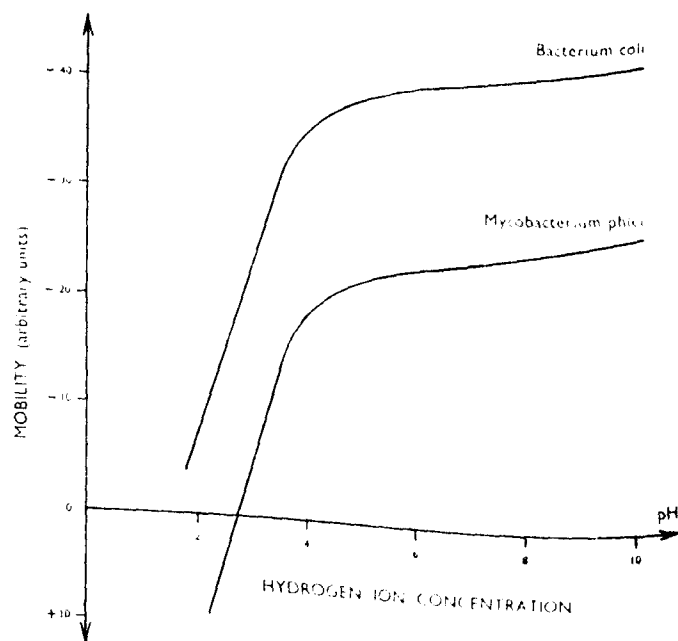


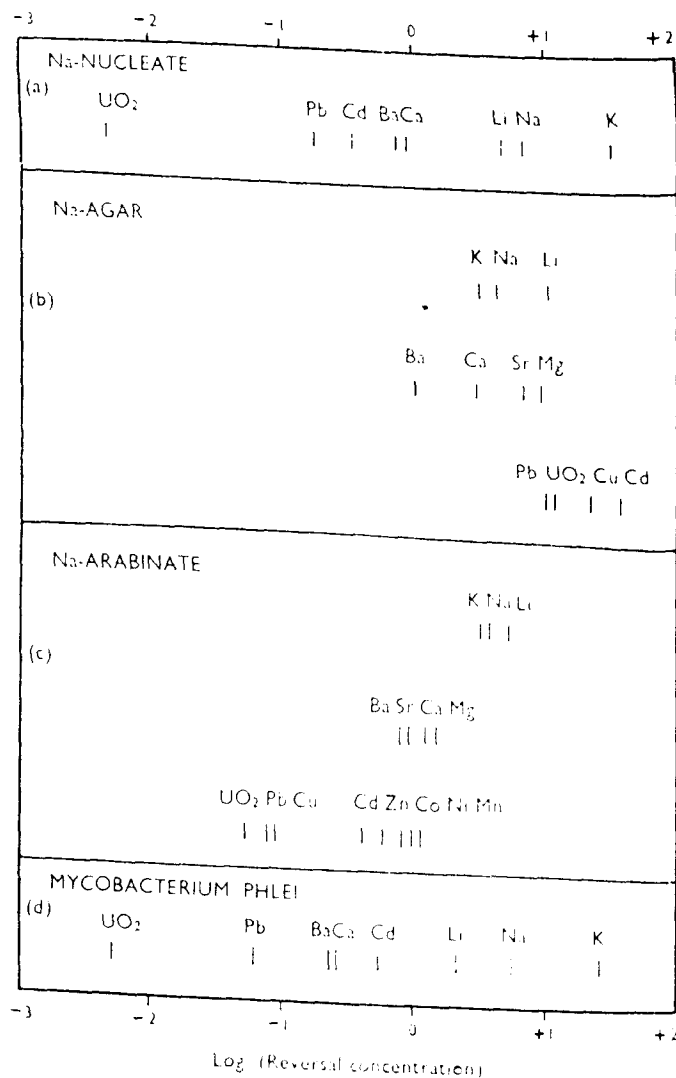
Fig. 1. Increasing the hydrogen ion concentration of bacterial suspension lowers the surface charge on the cells and may even reverse it.

negative charges within a molecule, or ion, can be displaced with respect to one another in an electric field which is, in the cases we are to consider, provided by an approaching ion. Considering the groups sulphate, phosphate, and carboxyl, which are the only negatively charged groups likely to be found in any quantity on a bacterial surface, it is believed that the polarizabilities follow the order $\text{PO}_4 > \text{COOH} > \text{H}_2\text{O} > \text{SO}_4$. The affinity between oppositely charged ions, the anion on the bacterial surface and the cation in the solution, depends on the valency, radius, and thus the polarizing power, of the cation, and on the polariz-

ability of the negatively charged group on the cell. It is also important which polarizability is the greater, that of water or that of the ionized group. Consider now the ions Li^+ , Na^+ , K^+ .

On the surface of these ions the fixation of a negatively charged ionized group will be easiest in the case of the smallest ion, Li^+ , and will be increasingly difficult for Na^+ and K^+ . Now, if the ionized group is more polar than water, the polarization energy will be added to the Coulomb energy owing to attraction of the opposite charges, and in this case the above order of cations will not be disturbed; in fact the spread may even be increased. The polarization energy for the smallest ion, Li^+ , after it has lost some of its water of hydration will be greater than for the others, and so, for an ionized group more polarizable than water, the order of decreasing affinity is Li^+ , Na^+ , K^+ . As these ions are gradually absorbed on to the cell surface the charge will be decreased and there will come a point at which it is zero. The concentration of any ion that causes this to happen for a given surface is known as the 'reversal of charge concentration'. Since reversal concentrations can be considered as an inverse measure of cation affinity, we must expect reversal concentrations to increase in the order Li^+ , Na^+ , K^+ . This is the order found for phosphate colloids. When the polarizability of the ionized group is less than that of water the polarization energy will oppose the Coulomb energy, and thus Li^+ will be fixed with the greatest difficulty and the order of affinity will be reversed.

Enough has now been said to show that by increasing the number of cations used to six or eight, preferably having a common anion to eliminate anion effects, the nature of a colloid surface may be decided with some accuracy. Using this approach the surface of *Bacterium coli* has been shown to be covered with carboxyl groups which originate in a mucopolysaccharide, and that of *Mycobacterium phlei*, an acid fast bacterium of the same class as the causative agents of tuberculosis and leprosy, has been shown to have a surface composed partly of a phosphatic fat and partly of protein. Figure 2 shows the reversal concentrations for this latter bacterium and their order is seen to be



similar to that of the phosphate colloid, which is only one of several phosphate colloids that have been characterized in this way. The case of *M. phlei* also involved the use of the anion effect. This acts in a manner similar to the reversal effects described with cations. If a surface contains positive groups, the mobilities in, say, M/100 NaCl and NaSCN will be different, but if the surface is composed of polysaccharide, for example, there will be little difference. In fact, there exist sequences of reversal concentrations with alkaloids by which many proteins may be distinguished, but nobody has yet tried to apply this to the study of bacterial surfaces.

Many variations even within one strain of bacteria are reflected in the electrophoretic properties of the cell. As far back as 1936 it was shown that the mobility of *B. coli* varied with age, and this must indicate changing proportions in some component of the surface. Furthermore, many bacteria exist in two modifications known from the appearance of their colonies on a solid medium as rough and smooth forms, and it was shown that in *B. coli* these have different mobilities. And to add even more interest to an already complex situation, it was recently reported that strains of this same organism, *B. coli*, exist which are sex incompatible and sex compatible and that they have different surface properties as shown by their different affinities for certain dyes. Thus, it can be seen that the surface charge upon a cell carries a great deal of information about several aspects of cell structure and function and that, if these charges can be interpreted, perhaps by the reversal of charge methods as much as by any other, then we shall be able to learn much about the outermost part of the living cell.

An interesting application of electrophoresis is the recent study of the training of *Aerobacter aerogenes* to the drugs pro-

Fig. 2. Reversal of charge concentrations of some common cations for (a) a phosphate colloid, (b) a sulphate colloid, (c) a carboxyl colloid, (d) *M. phlei*. The order of the reversal concentrations for *M. phlei* is seen to be similar to that of the phosphate colloid.

flavine and crystal violet. When this bacterium—and many other types of bacteria for that matter—is grown in gradually increasing concentrations of these drugs it survives and will eventually live in concentrations that would originally have been lethal to it. It is then said to be trained to this drug. This property is transmissible to subsequent generations even in the absence of the drug and it is sometimes accompanied by marked changes in surface properties. There is apparently more than one type of drug resistance because, considering only training to crystal violet, it is found that, depending upon the age of the parent culture from which the trained strain is produced, three different types of trained cell are borne. A young culture produces a so-called type *A*, which has exactly the same electrophoretic properties as the parent; an old culture produces a type *C* which has vastly different surface properties; and a culture of intermediate age, one with intermediate properties. This process will be mentioned again in the next section.

Infra-red studies

Almost all substances absorb radiation in the infra-red region of the spectrum, and by virtue of the different molecular vibrations possible in each molecule this is a highly characteristic property of a substance. The infra-red spectra of many substances of biological origin such as nucleic acids, mucopolysaccharides, fats and so on have been recorded. It was therefore a logical extension to take the spectra of bacteria. The first recorded instances came from America, but the work was taken up and put on a firm footing by workers in Great Britain. As a result it is now possible to identify bacteria by means of their infra-red spectra. Figure 3 shows the infra-red spectra of four species of bacteria. It will be noticed that for the most part the differences are small, but they are definite and what is more important they are repeatable. Provided that a carefully standardized technique is used, these small differences and, indeed, smaller ones, can be faithfully reproduced. A statistical method of comparison developed in Britain has been shown to be valid. So this puts the comparative study on a numerical basis rather

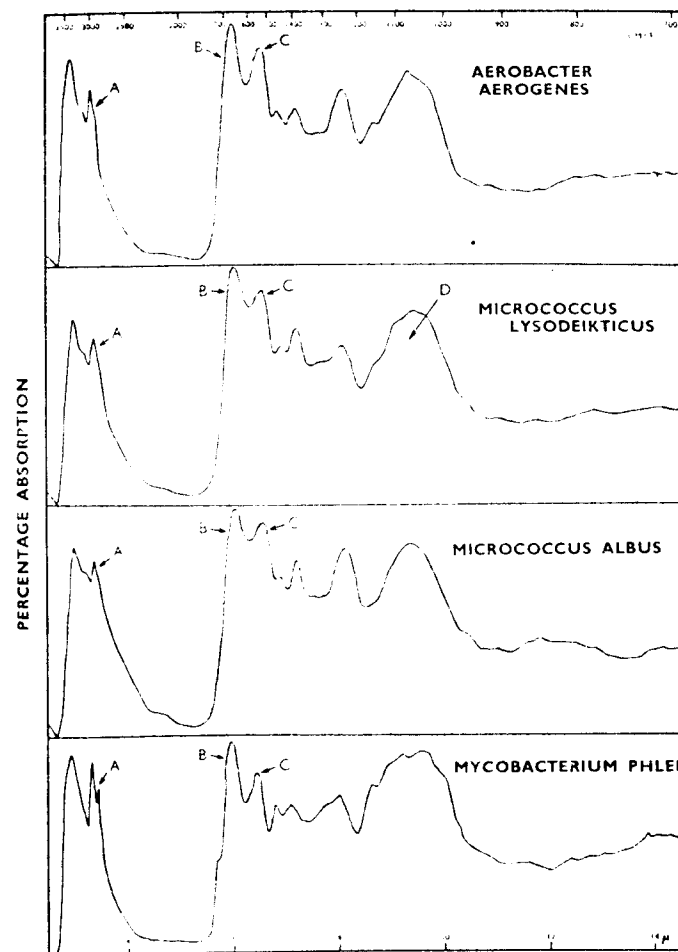


Fig. 3. Infra-red spectra of four different species of bacteria.

than a visual one involving direct comparison of the relative intensities of various peaks. By now spectra of some six or seven hundred different bacteria have been recorded, and apart from a very few closely related strains it is possible to distinguish between them all. This method is rapid and in time may to some extent replace the time-honoured methods of the medical bacteriologist. Interesting though this is in itself, it is clear that by virtue of the differences that mark them off from one another these spectra contain much information that is of interest to the cell biologist. Owing to the common structural features of each of the component fractions (e.g., proteins) almost all bacteria show absorption at certain common frequencies, but variation from species to species in the proportions of the components is the main reason for the variety of outlines exhibited by bacterial spectra. Attempts to interpret these spectra are still in their infancy, but a certain amount of progress has been made by using the following approach.

Comparison of the bacterial spectra with spectra of both simple and complex substances, including materials of biological origin, gives some indication of the origin of some of the absorption peaks. For example, the peaks marked *A* are always found in the spectra of compounds containing hydrocarbon units and are associated with vibrations of the C—H bond; the peaks marked *B*, *C* are always found to be intense in proteins. But to confirm such indications and to get more accurate information the chief method has been of selective extraction. Suppose we extract the fat of a cell with a simple fat solvent and leave the remainder of the cell intact. When we take the spectrum of the cell residue we find that the spectral contribution of the residue has been considerably increased with respect to its contribution to the original whole-cell spectrum. This is because it is necessary to take a constant weight of sample for each spectrum and we have thus magnified the contribution of part of the cell and made it easier to recognize. By using suitable chemical treatments this can be done for most of the different constituents of the cell in turn. It will be noticed that here we are studying the cell as a whole, but we can correct the results

to some extent for the contribution of the cytoplasm and so use it as a method for studying the cell wall. An alternative approach is to prepare cell walls as described above and use them.

At the moment it is not possible to give a complete interpretation of these spectra since even the spectrum of a protein cannot yet be fully explained, but the above approach can be used to provide a quick guide to the main components of the cell wall. For example, in *Micrococcus lysodeikticus*, which is a nasal commensal, treatment to remove polysaccharides caused the peak marked *D* to be considerably attenuated. Since there was no nucleic acid in the cell walls, and very little fat, and the remainder of the spectrum could to a large extent be explained by the presence of protein, it was concluded that the cell wall consisted of protein-polysaccharide mucilage. There was, of course, further evidence, but this is the type of approach in use and should give an indication of the information that can reasonably be expected from it.

A particularly interesting case is that of the training of *A. aerogenes* to crystal violet, mentioned above. The untrained and the type *C* strains had vastly different spectra although they bore a parent-progeny relationship and would be expected to have closely similar spectra. There had obviously been a profound alteration in the cell composition, and a fat was eventually isolated that had a spectrum containing all the unusual features

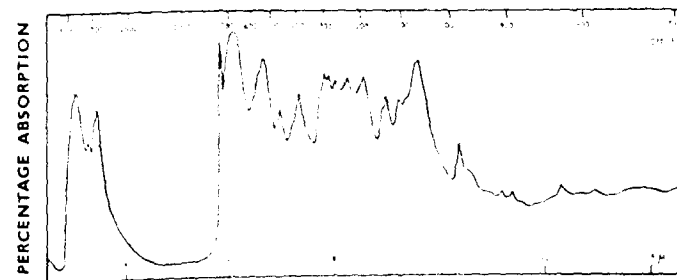


Fig. 4. Intra-red spectrum of *Aerobacter aerogenes*, type *C*, trained to crystal violet. This should be compared with that of the parent strain in Fig. 3.

in the spectrum of the type *C* cells. After removal of the fat the cell residue appeared to have a fairly normal spectrum, so the effect of training to the drug had been to alter the fat content of the cells. Not only had the concentration of fat been vastly increased but the fat itself was different from that produced by the parent before training. The full details of this training process are not yet understood and the scheme is no doubt complex since we have to account for three different trained strains. It is interesting to note that the spectrum of the type *A* trained strain is identical with that of the untrained parent. There are other applications of infra-red spectroscopy to microbiological problems, but space precludes their description here. Amongst the problems being studied by this method are the germination of bacterial spores, and further drug resistance phenomena.

Analytical Studies

Analysis of the bacterial cell wall by the methods of biochemistry is a valuable tool in the study of the components of the cell. The physical techniques described above help with the understanding of the arrangement of the cell, and to some extent with its composition, but to discover what amino-acids are present in the proteins of the cell, and what sugars form the polysaccharides, we must break down the parts of the cell into their ultimate units or molecules. The usual method is to break down the various parts, say the cell wall, by hydrolysis with dilute acid, and then identify the products chromatographically.

It is common knowledge that all bacteria can be placed in one of two classes according to whether they retain or lose the stain in a procedure known as the Gram reaction. The mechanism of this reaction is still unexplained, but nevertheless the division it makes seems to be a fundamental one. There is a marked difference in composition between the cell walls of Gram positive and Gram negative bacteria. The cell walls of Gram negative species contain about ten times more fat and ten times less glucosamine (an amino sugar) than do those of Gram positive cells, and their protein components exhibit almost the complete range of amino-acids found in proteins. On the other hand, the

cell walls of Gram positive species are found to contain much less fat and have a protein moiety which contains only three or four amino-acids, namely glutamic acid, alanine, either diaminopimelic acid or lysine, and sometimes glycine or serine. The distribution of the sugars varies greatly between the species, although glucosamine and another amino sugar, named muramic acid, are always present. These results point to a 'basal' cell wall structure built up of the following organic molecules: a hexosamine component that contains glucosamine, muramic acid, and sometimes galactosamine; a peptide component as described; and usually a polysaccharide containing not more than four sugar residues. This does not mean that the cell walls of these bacteria contain only these components. It is thought that this possible 'basal' structure of sugars, amino-sugars, and peptides which is pictured as linked up in a highly complex and unspecified manner, forms the basis upon which is attached the antigenic part of the cell wall. This latter part can be removed by proteolytic digestion and may contain amino-acids differing from those taking part in the 'basal' structure.

Gram negative species have not been investigated as intensively as Gram positive ones, and owing to the much larger number of constituents it may well take much longer to arrive at a full understanding of them. The Gram staining reaction is still unexplained, although it is apparent that the intact cell is essential for the reaction because parts of the cells are uniformly Gram negative no matter what their origin. The walls of Gram negative species, which may contain up to 25 per cent by weight of fat, are known to contain lipoprotein granules similar in size to those obtained by disintegration of the cytoplasmic membrane of Gram positive cells. But it is not yet known whether these granules are part of the cell wall proper, or whether they are part of the cytoplasmic membrane enmeshed in the cell wall. Under the electron microscope, walls from some Gram negative cells show two layers, the inner consisting of spherical macromolecules, a state that has not so far been reported amongst Gram positive species. It will be seen that since the cell walls of Gram negative species contain a large lipo-

protein fraction, there will be two contributions to the amino-acid analysis, namely that from the lipoprotein and that from any other proteinaceous component that may be there. Thus, the 'basal' structure found in Gram positive cells may well be present in Gram negative cells but may be screened from recognition.

Little work has been done on the constituents of the cell walls of other types of micro-organism, but there is reason to believe that the basal structure is confined to bacteria. The most important point about this is that it enables an explanation to be given of the selective action of penicillin upon bacterial cells. The varied effects observed when a bacterial culture is grown in the presence of penicillin, such as lysis and the appearance of distorted forms, can be explained if the action of penicillin is to inhibit one of the reactions leading to the synthesis of the basal wall structure. And there the situation stands: the existence of a basal wall structure in bacteria is proposed and is supported by much evidence. In time, further research will confirm or reject it, but the unifying nature of such an idea makes its reality seem eminently reasonable.

Kinetic studies

So far we have worked from the outside of the bacterial cell inwards: starting with the slime layer which can be studied by electrophoresis we proceeded to the structure and composition of the cell wall itself as revealed by infra-red spectroscopy and analysis, and we must now consider the cell as a whole, including its cytoplasm. One cannot, of course, completely separate the various parts of a system so intimately linked as the living cell, and it is when we come to consider the cell as a unit that this is most strongly felt. For some few years, research workers, notably those under Hinshelwood at Oxford, have been studying the kinetics of various processes connected with bacteria. Their work has done a great deal to remove some of the innate mystery surrounding biological reactions. Hinshelwood has demonstrated mathematically that reaction systems can be postulated which are self-propagating and behave in a similar manner to

many biological systems. 'It seems not unlikely that there are after all no individual units endowed with mysterious powers, but that the mystery, indeed the miracle, lies in the initial conjunction of partial systems which minister to one another's needs in the appropriate way.' Let us take an example from a very simple case of this sort, that in which two substances X and Y are being formed, each at a rate determined by the catalytic activity of the other so that

$$dX/dt = \lambda Y \text{ and } dY/dt = \mu X.$$

Neither X nor Y alone would increase in an autocatalytic manner, yet if we solve these equations we find that after a period of adjustment the amounts are given by

$$X = X'e^{kt} \text{ and } Y = Y'e^{kt}$$

where X' , Y' are values at $t=0$. It can further be shown that even when more complex interdependent systems are postulated involving several stages and bifurcations, the final law governing the formation of the various products is one of exponential increase. These systems therefore appear to be autotrophic.

Turning to the practical side of the problem, it is found that there is much evidence to support this type of mutually dependent reaction pattern. For example, it is well known that the amount of nucleic acid in bacteria at any one time is related to the rate of growth at that stage and that the equations

$$d(\text{Protein})/dt = \lambda (\text{Nucleic acid})$$

and

$$d(\text{Nucleic acid})/dt = \mu (\text{Protein}),$$

fit the observed quantities fairly closely. Selective disturbance of the workings of the cell is possible by growing the cells in the presence of drugs and chemicals that, for one reason or another, block a certain reaction and cause some apparent upset of cell economy. This may manifest itself by the appearance of an unduly long lag phase, the phase that precedes the period of active growth and division known as the logarithmic phase, or the appearance of filamentous and deformed cells in the culture. By means of these and other carefully designed experiments evidence is slowly accumulating to show how the cell works,

how the cell wall is built up, how the cell prepares for division, and the means whereby unusual metabolites are incorporated into the synthetic scheme and lethal drugs dealt with.

Conclusion

There are so many types of bacteria, and their activities cover such an immense range, that it will be a long time before their activities and properties are understood in any completeness. We are now beginning to understand how the cell is made up, in what manner the various parts are arranged, and what functions they have. It is indeed only an outline, but a good beginning has been made and in filling in some of the details the methods described above will probably play an important part. They are not the only methods used to study bacterial problems but they are amongst the more valuable from the point of view of understanding structural composition.

FURTHER READING

Work, E.: 'Biochemistry of the Bacterial Cell Wall', *Nature* 1957 **179** 841.

Hinshelwood, C. N.: *The Chemical Kinetics of the Bacterial Cell* (Oxford, 1946).

There is no review at present available covering the electrophoretic and infra-red work, but the following papers will give more information:

Thomas, L. C., and Greenstreet, L. E. S.: *Spectrochim. Acta* 1954 **6** 302.

Barry, P. J., and James, A. M.: *Biochem. et Biophys. Acta* 1954 **15** 186.

Lowick, J. H. B., and James, A. M.: *Biochem. et Biophys. Acta* 1956 **17** 424.

Davies, J. T., Haydon, D. A., and Rideal, E. K.: *Proc. Roy. Soc.* 1956 **B145** 375.

Rideal, E. K., Adams, D. M.: *Chemistry and Industry*, 1957 762.

MINERAL DRESSING

E. J. PRYOR

IN 1952 the University of London instituted a three-year course leading to a B.Sc in 'mineral dressing'. Other countries are now doing something similar but the demand for scientists trained for this important work is far from met. The purpose of this article therefore is to explain what the term 'mineral dressing' means – to explain how it fits into the pattern of modern technology.

The minerals that justify mining and working up into metals or other basic materials in industrial use, are mostly found in the form of low-grade ore deposits. In the past the rich ores were often spectacular, with easily visible crystals or nuggets set in lodes, lenses, or other deposits from which they were easily separated by a few manipulations based on crushing and the use of gravity. Today the ores mined are usually very low in assay grade, and the valuable mineral is often disseminated through the barren portion of the lode, or vein, in the form of specks so small that they are hard to see without a microscope. To smelt the whole of such run-of-mine ore in order to extract its one or two per cent of tin, copper, or other metal would be an expensive and wasteful business. The first stage in extraction metallurgy is therefore to remove as much of the valueless rock as cheaply as may be, in order that the retained fraction can be treated economically by more costly methods. This upgrading is the work done by the mineral dresser, who takes the ore as it comes from underground, liberates the valuable specks from their matrix and then sorts out these specks. In this operation, methods are used that employ many of the resources of physics and chemistry, in operations at least as involved as those used in chemical engineering.

One big difference between mineral dressing and a chemical

process is that the raw materials in mineral dressing are stony and almost inert. Another is that in most of the processes used the physical structure of the ground particles – whether valuable concentrate or barren tailings – remains unaltered, save for minute changes induced by treatment with surface-active agents.

The work done in the mill is not confined to the needs of metallurgy. The phosphates used by the farmer, special cements, abrasive garnets, uranium for atomic piles, asbestos for car brakes, talc for face powder, china clay for fine ceramics, low-ash coal, diamonds, fluor spar – all these are among the grist ground from the huge tonnages of rock treated every day by one or other process.

Treatment starts as near the mine as the lie of the country permits. The first stage is the liberation of the valuable constituents. The big lumps of rock from the mine go through crushers to bring them down by stages to a diameter of an inch or less. Since modern mining is heavily mechanized, waste rock is almost always mixed with the true ore sent to the mill, but during this dry-crushing it is often possible to remove some of it. Wherever the colour or lustre of the ore is sufficiently marked to be readily observed as the coarse run-of-mine rock moves along the conveyor belts on which it is displayed, hand-sorting is still used. If ore is radio-active, a Geiger Muller counter can take over, each piece of radioactive ore signalling its presence and so causing an automatic deflector to remove it from the passing stream. If the waste has a density of, say, 2.7 to 3.0, and the value-bearing ore is 3.1 or more, these lumps can be separated at this early stage by the dense-media process. In this, the ore and waste are fed into slurry of ferro-silicon and water maintained at a splitting density of 3.05. The heavier pieces sink and the light ones float to their separate discharge gates, thus leaving less of the original tonnage to go on to further treatment.

After this preliminary sorting and crushing the ore is mixed with water and runs the gauntlet through horizontal cylinders half filled with steel rods or balls, which churn and batter it to a fine pulp. The process is continuous and accurately controlled, its object being to free each valuable speck of mineral fairly

completely, and to end the grinding process just as soon as this has been achieved. Some of the physical and chemical processes used in the subsequent stages of treatment are very delicate, and will not work properly if the particles are so big as to have a fast settling rate in water. It is a measure of the accuracy possible in mineral dressing that even particles only 2 or 3 μ in diameter can be successfully separated if their value justifies the cost, even when hundreds of tons are treated every hour at these microscopic particle sizes.

The methods of mineral dressing most in use before World War I were limited to gravitation in pulsing or streaming currents of water (*jigging*, *sluicing*, and *tabling*), the use of magnets on ferro-magnetic ores, and such hydro-metallurgical processes as the cyanidation of gold and the leaching of copper.

A particle drifting down through still water reaches a terminal speed in accordance with the force of gravity acting on it and of the retarding effect of the laminar flow of water past its rubbing surface. If the particle is big, and one of a crowded layer, it can be tossed up and down by violent pulsation till it stratifies according to mass and shape, and the strata can be separated continuously. This is the process of *jigging*. In *tabling*, used for fine sands, the crushed ore is rolled downstream in a nearly horizontal device to which skidding reciprocations are applied to produce a separating effect. If the particle is very small its mass-effect can be increased by the use of centrifugal force in the separating machine.

The overwhelmingly important process in use today is based on differences in the surface tensions of the various minerals liberated when a lump of ore is ground to the fineness of flour. The development of this powerful technique has been swift, with practice often ahead of our knowledge of principles. The process depends on a difference between the surface tension of one mineral in the pulp and that of all the other kinds. This tension is developed so as to render the surface of one mineral species aerophilic (air-loving) while the rest remain swimming. Air is then mixed with the pulp. The aerophilic particles attach themselves to air bubbles and rise to form a froth which can be

removed with its valuable load while the unwanted ones swim away down the tailings race. The particles floated (as those sticking to the air-bubbles are described) must not be so heavy that their gravity-pull can overcome the holding effect of their grip of the air.

Now let us consider a typical differential flotation, in which the object is to separate first galena (lead sulphide), next sphalerite (zinc sulphide), and finally pyrite (iron sulphide) from the ore, leaving only a barren siliceous tailing. First the pH is adjusted to an alkalinity of between pH 8 and 9, using, in the grinding mill, sodium carbonate, with perhaps a little sodium silicate, and a trace of cyanide to 'kill' any random metal ions. The silicate is a 'dispersing agent' added to keep the newly sheared surface of the particles free from slime, which otherwise could mask these surfaces and make any further 'conditioning' impossible. This 'conditioning' is a process in which inherent differences between the surface reactions of the various mineral species are developed. A very small quantity of collecting chemical (usually potassium ethyl xanthate) is stirred in. Its molecules are attracted to the positively charged lattice points of the ground sulphide minerals, which, being newly sheared, are in a highly reactive state.

When a crystal structure is broken, the atoms forming the new surface are partially reft from the other atoms around them. A crystal structure is a regular repetition of a system of bonds between its component atoms, arranged to satisfy the energetics of the closed system at the lowest level of excitement. The effect of a sudden break is to leave the surface atoms unbalanced and electrically charged at the sites where these atoms are built into the crystal – the lattice sites. In mineral grinding this excited surface, with its positively and negatively charged lattice points, is suddenly formed in the aqueous pulp. At this moment, if the mineral dresser has done his job correctly, a molecule of xanthate swims up from the water surrounding the particle, sheds its potassium ion and hooks on to a positively charged lattice point by means of its now negatively charged group, while the ethyl alcohol, which forms the rest of the xanthate

molecule, waves above the scene like an anchored balloon. That is the picture we have of the process at our present state of research.

If indiscriminate coupling of this sort took place, the process would fail, so it is necessary to add other chemicals (mere traces, of the order of 0.1 lb/ton of ore, or less). Zinc sulphate is used to produce a common-ion loading of the interphase or zeta zone round the sphalerite, while the pyrite usually remains inactive so long as the pH exceeds 6. The silica attracts silicates and hydroxyl groups, and retains a hydrated surface throughout the treatment.

At this stage the galena particles have become aerophilic while all the other minerals are hydrophilic. The pulp is now subjected to intense aeration in the presence of a trace of cresylic acid or of a terpene. The galena is picked up by air bubbles and floats to the surface of the aerating tank where it becomes part of a mineralized froth, sufficiently stabilized by the terpene to be skimmed off before bubbles explode, and a rich concentrate of minute particles of galena is recovered.

Next comes the sphalerite. From the first line of flotation cells – the tanks in which the aeration is carried out – the pulp flows to another conditioning treatment. The pH is checked and raised to 9.5 or so with lime. A trace of copper sulphate is added to the pulp, which is stirred for several minutes. This gives time for copper ions to leave the aqueous phase and deposit on the zinc lattice points at the surface of the sphalerite particles. This almost undetectable surface change makes practically no difference to the essential composition of the zinc sulphide, but from now on the copper envelope ensures that, so to speak, its zinc contents are delivered at the correct address. More xanthate is added, and the pulp goes forward to another row of flotation cells from which a new mineralized froth is removed, this time of high-grade zinc sulphide.

The tailing from these cells is now adjusted with sulphuric acid to a pH of 5 or 6 and again 'floated' after a further addition of xanthate. This time the pyrite rises in the froth and the now valueless silica is run to waste.

This is but one of the many treatments going on at the present time in thousands of mills day and night. Many variables affect the course of the work, originating from slight changes in the toughness of the ore, in changes of purity in the local water as rainfall varies seasonally and affects the salts carried in solution, and other causes. Since it is important to any flowline to have a smooth rate of travel through each processing stage, automation is used wherever possible to operate the controls. While the ore is still lumpy, any power failure or breakdown along the transporting, crushing, and screening lines must be dealt with promptly, otherwise hundreds of tons of rock may pile up and damage the system. Closed-circuit television cameras scan such lines, and photo-electric cells warn the operator when a conveyor belt is running empty or a holding tank is too full. At key points the pH of the pulp is checked and recorded, and electrically operated valves adjust the entry of water and chemicals in accordance with changes in the composition of the pulp so as to maintain a correct ratio between solid and liquid at the desired alkalinity.

The scientists who make the original tests and pilot-scale runs on a given ore have a fascinating job. All research work contains elements of detection, but as a rule they are watered down by a monotony of measuring, weighing, and titrating on tedious lines. There is more of the thrill of a good 'whodunit' for the mineral dressing research worker, as is apparent when his apparatus is considered. He needs microscopes that go down to a micron or less, and micro-analytical apparatus to pick out a minute grain and analyse it in a laboratory which, though comprehensive, would pack into his brief-case. He must often use apparatus for taking powder X-ray pictures, for spectrography (emission and sorption) and for interpreting results yielded by this kind of probing. Then there are elutriating devices which will sort out particles in liquids or air down to a few microns in diameter. The size variable brought under control, the possibilities of separation are checked in terms of density and surface differences. From this work emerges a promising method of treatment.

Next comes the pilot plant, in which the isolated or batch tests are run together in continuous form, and modified until a steady treatment is possible. It must be practical. Too sensitive a process may give too much trouble to be risked in the field, however prettily it works in the laboratory. And so by logical steps we go to the stage where the engineers take over and the ton-an-hour pilot plant expands to a 5,000-ton-a-day mill. The end-product may be copper for electrical installations, the niobium in the special alloy of a plane's jets, or the vanadium steel in the stub-axle that makes a car safer to steer. Mineral dressing is an essential step in bringing most minerals to the service of men, whether in the novel use of uranium to produce nuclear power or the ancient one of getting the tin out of a Cornish mine.

These examples show the wide range of problems the mineral dresser has to solve, and they are only part of the story. In the university course where specialized training is now given, the undergraduate is first trained in fundamentals of chemistry, physics, and geology. He then specializes in microscopy, surface chemistry, metallurgy, and engineering to the point where he can recognize through his specialization the principles of his science, inspect them in action in the laboratory, and finally control them on the large scale called for in the running of ore treatment plant. His job may take him all over the world. It combines several sciences so that he never ceases to learn, and never loses interest when, with increasing experience, his responsibilities grow. What, perhaps, is the most rewarding feature of this career is that those who follow it are meeting a need without which civilization could neither advance its standards of living nor even maintain them, since without the flow of more and more metals, rare earths, and other refined mineral products, into the basic industries, man would find it even more difficult to meet the alarming rise in the birth-rate which has already taken place in many lands, let alone push back the waste lands and tame the barren places for his use.

FURTHER READING

- Richards and Locke: *Text Book of Ore Dressing* (McGraw-Hill 1940).
 Taggart, A.: *Handbook of Mineral Dressing* (Wiley, 1945).
 Pryor, E. J.: *Introduction to Mineral Dressing* (Mining Publications Ltd., 1955).
 Rabone, P.: *Flotation Plant Practice* (Mining Publications Ltd. 1957).
 Sutherland and Wark: *Principles of Flotation* (Australian I.M.M., 1955).
 Gaudin, A. M.: *Flotation* (McGraw-Hill, 1957).

THE ORBIT OF A SATELLITE

W. C. VAUGHAN

It requires little imagination to recognize in the American and Russian projects for launching earth satellites during the International Geophysical Year enterprises that may be as epoch-making as any in man's history. Success in these ventures at once poses new challenges to the ingenuity of the engineer. Not the least of these, because it is generally regarded as offering the most promising means for ultimate escape from the vicinity of the earth, is the possibility of modifying the orbit of the satellite. Much has already been written on the subject of earth satellites, but current literature has tended to fall into two distinct categories. The first of these is intended for readers who have a good grasp of the essentials of celestial mechanics; the second, descriptive and often highly imaginative, makes little demand either on the reader's acquaintance with the laws of motion, or on any ability to use such knowledge of elementary mechanics as he may possess. There must be many, however, who although they desire a more exact understanding of this fascinating subject have been led to believe this to be beyond their attainment. Textbooks on mechanics do not usually make very attractive reading, and they are not infrequently exasperating when, in their treatment of planetary motion, they refer the reader to other works for some proposition associated with the geometry of conic sections, or to justify some step in calculus. It is, indeed, not generally realized that a very clear appreciation of the main principles of the subject can be obtained from a rudimentary knowledge of Newton's Laws, backed by very modest mathematical ability. The purpose of this article is to show how this can be achieved.

Circular Orbits

It is well known that any body falling freely towards the earth

does so with an acceleration inversely proportional to the square of its distance from the centre of the earth. This acceleration is, of course, the manifestation of the gravitational attraction between the body and the earth, and it is directed towards the earth's centre. Near the surface of the earth, i.e., at about 4,000 miles from the centre, a falling body which is negligibly impeded by motion through the atmosphere gains velocity at a rate of about 32 feet per second in each second of fall. This acceleration is usually denoted by the symbol g . For calculations on a celestial scale, it is preferable to quote this as the gain in velocity in miles per second in each second; g then has a value of $32/5280$ or 0.00606 miles per second per second. At a distance of, say, r miles from the centre of the earth, a freely falling body will have an acceleration of

$$\begin{aligned} a &= 0.00606 \times (4,000)^2 / r^2 \\ &= 97,000 / r^2 \text{ miles per second per second.} \end{aligned}$$

According to the first of Newton's Laws of Motion, a body moving in a straight line through space will continue to do so with unchanging velocity unless it is influenced by an external force. If it is to be made to move in a circle with constant velocity, then it must continuously experience a constant attractive force directed towards the centre O of the circle. Now a constant force produces a constant acceleration in the direction of the force. It is fairly easy to show that to maintain a constant velocity v in a circle of radius r , the body must suffer an acceleration v^2/r towards O .

Let us combine the information contained in the two preceding paragraphs. A satellite circling the earth with constant velocity v miles per second at a constant distance r miles from the earth's centre O experiences an acceleration v^2/r miles per second per second towards O . An acceleration in the required direction equal to $97,000/r^2$ miles per second per second is, however, provided by gravitational attraction. This will suffice to maintain the satellite in its orbit if

$$\frac{v^2}{r} = \frac{97,000}{r^2}$$

that is, if the orbital velocity is

$$v = 311 / \sqrt{r} \text{ miles per second.}$$

Here, then, we have a simple relation from which it is possible to calculate, in terms of its height above the earth, the orbital velocity of any satellite that moves in a circular path about the earth. Thus, for example, a satellite circling at, say, 900 miles above the earth is moving in an orbit of radius 4,900 miles. To maintain its path, it must have an orbital velocity of $311/70$ or roughly 4 miles per second.

It is also an easy matter to calculate the time required to complete one circuit of the earth. The circumference of the orbit is $2\pi r$ miles and the orbital velocity is $311/\sqrt{r}$ miles per second. The time T to complete one circuit, known as the period of the satellite, is therefore given by the relation

$$\begin{aligned} T &= 2\pi r \times \sqrt{r}/311 \\ &= 0.02 \times \sqrt{r^3} \text{ seconds.} \end{aligned}$$

For a satellite orbiting at 900 miles above the earth's surface, the periodic time is therefore about 6,900 seconds, or 115 minutes. Figure 1 shows the orbital velocities and periodic times for satellites in circular orbits for heights up to 1,500 miles. The foregoing calculations assume, of course, motion in an atmosphere so tenuous as to have no appreciable effect. Atmospheric drag would tend to slow down the satellite and reduce its energy of motion. This would lead to shrinkage of the orbit, the fall towards the earth restoring some of the lost energy of motion. Approach towards the earth would, however, bring the satellite into a denser atmosphere: the drag would increase and the orbit would continue to shrink.

Modification of the Orbit

It is evident that in order to launch a satellite into a circular orbit of given radius r , two very exacting conditions must be satisfied. The velocity of the satellite must be $311/\sqrt{r}$ miles per second, and this must be in a direction exactly at right angles to the line from the satellite to the earth's centre. If either of these conditions is not met, the intended course will not be followed.

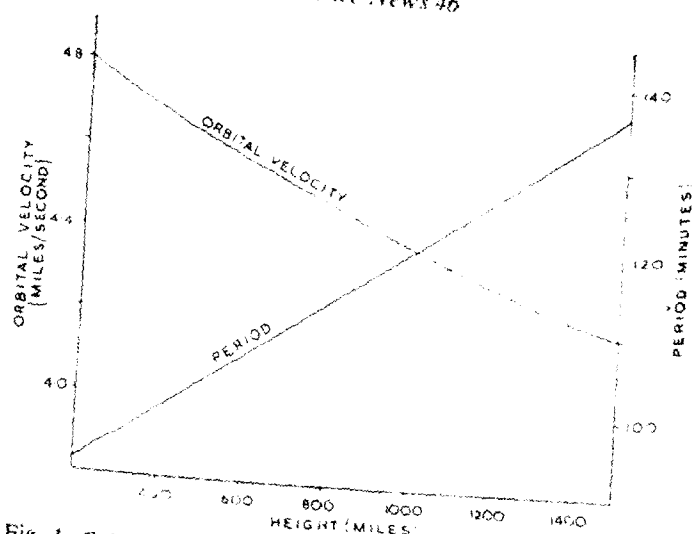


Fig. 1. Orbital velocity and period of a satellite in a circular orbit. This figure may also be used for satellites in elliptical orbits. If the horizontal scale represents the perigee distance, the left-hand vertical scale multiplied by $\sqrt{1+e}$ gives the perigee velocity; the right-hand vertical scale divided by $\sqrt{1-e}$ gives the period. If the horizontal scale represents the apogee distance, the left-hand scale multiplied by $\sqrt{1-e}$ gives the apogee velocity; the right-hand scale divided by $\sqrt{1+e}$ gives the period.

Conversely, if successfully placed in the required orbit, any subsequent modification of the velocity of the satellite will cause it to depart from its original path. In general, the new orbit will not be a circle. To appreciate the consequences of such an alteration in velocity, let us consider a particular case. It will be seen from Figure 1 that a satellite at a constant height of 1,000 miles above the earth, i.e., in a circular orbit of radius 5,000 miles, will have an orbital velocity of about 4.4 miles per second. Suppose that in some way it can be given an additional outward velocity of 0.8 miles per second as indicated in Figure 2. What then will be its subsequent path?

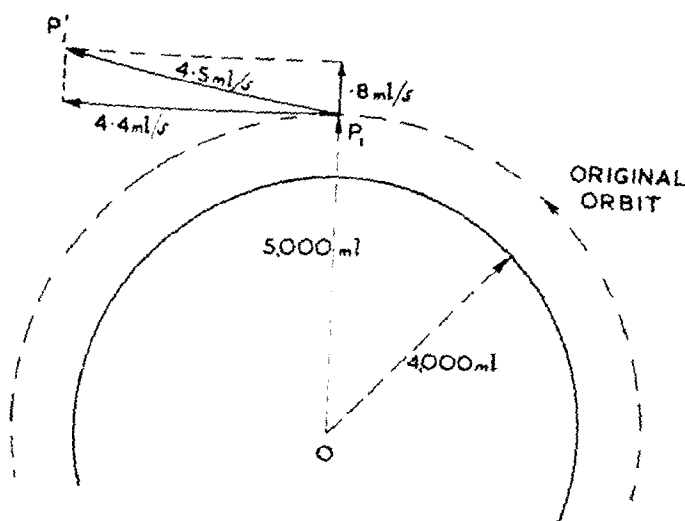


Fig. 2. Displacement of a satellite from a circular orbit. The satellite, originally in a circular orbit of radius 5,000 miles with an orbital velocity of 4.4 miles per second, is displaced by giving it an additional outward velocity of 0.8 miles per second. It will commence to move in the direction P_1P_1' with a velocity of 4.5 miles per second.

To solve this problem we shall employ a piece of information derived from the well-known law of celestial mechanics which was first propounded by Kepler early in the 17th century. This states that when a body moves under the action of a central force, i.e., a force directed always towards a fixed point, equal areas are swept out in equal intervals of time by the line from the body to the fixed point. A satellite is such a body, provided that it is outside the atmosphere, yet sufficiently close to the earth for it to be regarded as uninfluenced by any other body. The centre of the earth is, in such cases, the fixed point.

Suppose that the path of the body is represented by the curve shown in Figure 3, and that the same time t is taken to travel

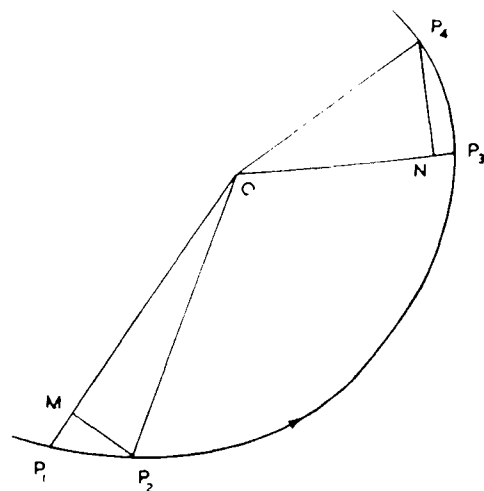


Fig. 3. *Kepler's second law.* A satellite moving in any orbit under the influence of a force directed towards O does so in such a way that the line from it to O sweeps out equal areas in equal times. If the time taken to travel from P_1 to P_2 is the same as that taken to travel from P_3 to P_4 , then the sectors OP_1P_2 and OP_3P_4 have equal areas. This means that the distance OP_1 multiplied by the velocity in the direction MP_2 must equal the distance OP_3 multiplied by the velocity in the direction NP_4 .

from P_1 to P_2 as from P_3 to P_4 . If O is the centre of attraction, then according to Kepler's Law, the sectors OP_1P_2 and OP_3P_4 have equal areas. If we make the time interval t sufficiently short, then the areas swept out are triangles. Their areas are $\frac{1}{2}OP_1 \times P_2M$ and $\frac{1}{2}OP_3 \times P_4N$ respectively, where P_2M is perpendicular to OP_1 , and P_4N is perpendicular to OP_3 . Now P_2M is the distance travelled by the satellite in the time t in a direction at right angles to OP_1 . If therefore v_1 is its component of orbital velocity in the direction at right angles to OP_1 , $P_2M = v_1t$. Similarly $P_4N = v_2t$, where v_2 is the component of orbital

velocity at right angles to OP_3 . If OP_1 is a distance r_1 and OP_3 a distance r_2 , the law of equal areas becomes

$$\frac{1}{2}r_1v_1t = \frac{1}{2}r_2v_2t.$$

This leads to the important conclusion,

$$r_1v_1 = r_2v_2.$$

It follows accordingly that, in any given orbit, the velocity of the body measured at right angles to the line joining it to O , multiplied by its distance from O , is always the same. This is a fact of particular significance in delineating the orbit of a satellite.

The second important principle that we shall apply in the problem we are examining is that the velocity of the satellite at any instant after displacement from its circular orbit is the resultant of its original orbital velocity and the velocity which brought about the displacement. In the position P_1 (Figure 2), the satellite has a velocity of 4.4 miles per second at right angles to OP_1 (its original orbital velocity), and 0.8 miles per second outwards along OP_1 (the displacing velocity). From the diagram it will be seen that these have a resultant of about 4.5 miles per second directed along P_1P_1' . The satellite will commence to move in this direction, but is prevented from following a straight line by the attraction towards O .

Suppose that after a certain time the satellite has reached a point corresponding to P_2 in Figure 4(a). As before, its motion at this instant is the resultant of a velocity of 4.4 miles per second at right angles to OP_2 (the original orbital velocity), and a velocity of 0.8 miles per second in a direction parallel to OP_1 (the displacing velocity). The diagram shows the resultant to be a velocity of about 3.8 miles per second directed along P_2P_2' . We therefore know the orbital velocity and the direction of motion at this instant. To determine how far P_2 is from O , we need to know the velocity of the satellite in the direction at right angles to OP_2 . This also can be easily obtained from the two components which contribute to the orbital velocity. Figure 4(b) shows that the displacing velocity is equivalent to 0.7 miles per second at right angles to OP_2 in a clockwise direction, and 0.4 miles per second out along OP_2 . The resultant

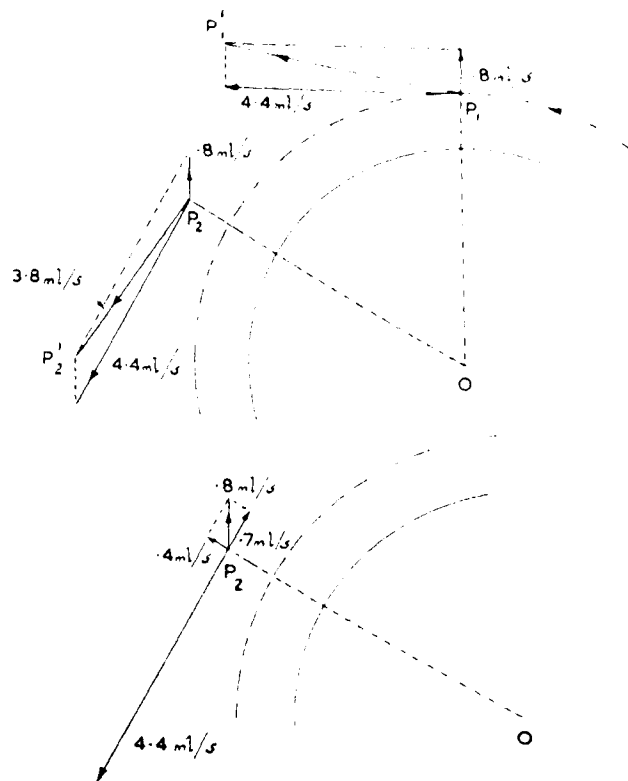


Fig. 4. Displacement from a circular orbit. (a) At P_2 , the satellite has the same components of velocity as it had at P_1 , namely, a velocity of 4.4 miles per second perpendicular to the line from it to O , and the displacing velocity of 0.8 miles per second in a direction parallel to OP_1 . The satellite will therefore move in the direction P_2P_2' with a velocity of 3.8 miles per second. (b) To apply Kepler's Second Law, we require the velocity at right angles to OP_2 . The displacing velocity is equivalent to 0.7 miles per second at right angles to OP_2 , and 0.4 miles per second along OP_2 . The resultant velocity of the satellite at right angles to OP_2 is therefore 3.7 miles per second.

velocity at right angles to OP_2 in a counter-clockwise direction is thus $(4.4 - 0.7)$ or 3.7 miles per second. According to Kepler's Law, this multiplied by the distance OP_2 must equal the corresponding product at P_1 , i.e., $3.7 \times OP_2 = 4.4 \times OP_1$. OP_1 is, however, known to be 5,000 miles, so that OP_2 must be 5,946 miles.

If this process is repeated, the positions of any number of points on the orbit can be found, and it quickly becomes evident that the latter is an ellipse. Of particular interest are the points P_A and P_P in the diagram of the orbit in Figure 5. These are known respectively as the apogee and perigee of the orbit; their positions are particularly easy to locate. At apogee, the original

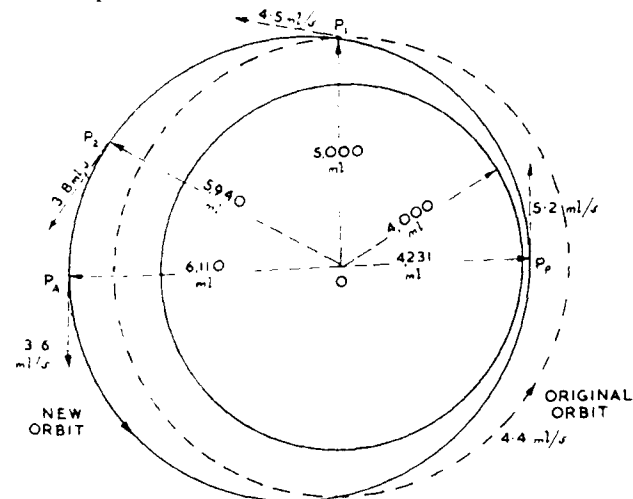


Fig. 5. The shape and dimensions of the new orbit. At apogee P_A , the displacing velocity of 0.8 miles per second is in a direction exactly opposite to that of the original orbital velocity of 4.4 miles per second. The apogee velocity is therefore 3.6 miles per second. At perigee P_P these two velocity components are in the same direction. Thus the perigee velocity is 5.2 miles per second. Note that at both apogee and perigee the orbital velocity is in a direction at right angles to the line from the satellite to the centre of attraction O .

orbital velocity and the displacing velocity are in exactly opposite directions, so that the orbital velocity at this point is $(4.4 - 0.8)$ or 3.6 miles per second. This is clearly at right angles to OP_A so that $3.6 \times OP_A = 4.4 \times OP_P$. The apogee distance OP_A is therefore $(4.4 \times 5,000) / 3.6$ or 6,111 miles. Similarly at perigee, the orbital velocity is $(4.4 + 0.8)$ or 5.2 miles per second in a direction at right angles to OP_P . The perigee distance OP_P is therefore $(4.4 \times 5,000) / 5.2$ or 4,231 miles.

The distance between P_A and P_P , the major axis of the orbit, is 10,342 miles. The minor axis is the width of the orbit measured at right angles to P_AP_P at its mid-point. This, from Figure 5, is 10,180 miles. The area of the ellipse is

$$A = \pi \times \frac{\text{major axis}}{2} \times \frac{\text{minor axis}}{2}$$

In this problem, therefore,

$$A = \pi \times \frac{10,342}{2} \times \frac{10,180}{2} = 82,680,000 \text{ square miles.}$$

We require this result in order to calculate the time taken for one circuit of the orbit. At apogee, as we have seen, the satellite travels a distance of 3.6 miles in one second in a direction at right angles to OP_A . The area swept out in one second by OP_A is therefore a triangle of base 3.6 miles and height 6,111 miles, i.e., about 11,000 square miles. According to Kepler's Law it will continue to sweep out area at this rate throughout the orbit. Thus the time necessary to complete one revolution must be

$$T = \frac{82,680,000}{11,000} \text{ seconds.}$$

This gives a period of about 125 minutes. In this time the satellite recedes to a distance of just over 2,110 miles from the earth's surface, and then approaches to within about 230 miles.

Elliptical Orbits

An ellipse is the path traced out by a point which moves so that the sum of its distances from two fixed points O and O' is always the same. A knowledge of the main dimensional features

The Orbit of a Satellite

5.

of this figure will enable us to translate the analysis of the particular problem we have just considered into more general terms. Figure 6 shows an ellipse having a major axis AA' of

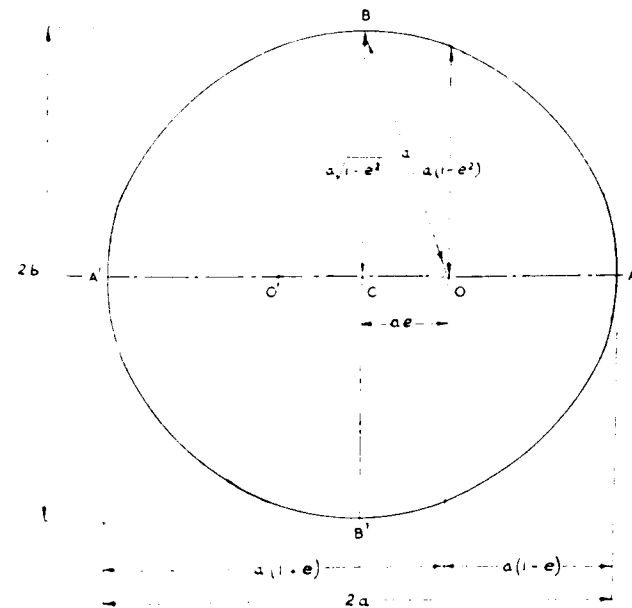


Fig. 6. The geometry of the ellipse.

Major axis $AA' = 2a$

Minor axis $BB' = 2b$

Eccentricity $e = \frac{CO}{CA}$

$OB = b$

$CB = b = a\sqrt{1-e^2}$

Area $= \pi ab$

$= \pi a^2 \sqrt{1-e^2}$

Apogee distance $OA' = \frac{(1+e)}{(1-e)} \times \text{perigee distance } OA$

length $2a$, and a minor axis BB' of length $2b$. The ratio CO/CA is termed the eccentricity e . From the definition of the figure, it follows that OB must be equal to CA . The main dimensions of the ellipse are now evident.

A satellite experiencing gravitational attraction towards O will, as has been shown, require a velocity $v_0 = 311/\sqrt{r_0}$ miles per second at right angles to OP if it is to maintain a circular orbit of radius r_0 miles. If it is displaced by giving it an outward velocity v_d along the direction OP , we may expect it to take up an elliptical orbit. Eventually it will arrive at apogee A' with a velocity $(v_0 - v_d)$ at right angles to OA' . Later it will reach perigee at A with a velocity $(v_0 + v_d)$, again at right angles to OA . According to Kepler's Law $(v_0 - v_d) \times OA'$ must equal $(v_0 + v_d) \times OA$.

$$\text{Thus } \frac{OA'}{OA} = \frac{v_0 + v_d}{v_0 - v_d} \text{ or } \frac{1 + v_d/v_0}{1 - v_d/v_0}.$$

Figure 6 shows, however, that OA'/OA is equal $(1+e)/(1-e)$. It is evident, therefore, that e must equal v_d/v_0 , i.e., the eccentricity of the orbit is the ratio of the displacing velocity to the original velocity. The eccentricity of the orbit in the numerical example considered in the earlier paragraphs was therefore $2/11$.

It is perhaps rather more realistic to consider the orbit in relation to the distance and velocity of the satellite at perigee. If these are respectively r_p miles and v_p miles per second, then $r_p = r_0/(1+e)$ and $v_p = v_0(1+e)$. From the relation $v_0 = 311/\sqrt{r_0}$ miles per second, it is not difficult to show that

$$v_p = \frac{311}{\sqrt{r_p}} \times \sqrt{1+e}.$$

Thus, if we regard the horizontal scale of Figure 1 as representing the perigee distance of a satellite in an elliptical orbit of eccentricity e , the perigee velocity may be found by multiplying the left-hand vertical scale by the factor $\sqrt{1+e}$. In a similar manner, if the horizontal scale is regarded as the apogee distance, then the apogee velocity is obtained by multiplying the vertical scale by $\sqrt{1-e}$.

A number of problems relating to the launching of earth satellites can be examined by means of the information now at our disposal. Suppose, for example, that it is intended to place such a body in a circular orbit at a height of 300 miles above the earth. Figure 1 shows that this requires that the satellite be given a velocity of 4.74 miles per second in a direction at right angles to the line from the satellite to the earth's centre when at the stipulated distance. If, in fact, it attains a rather smaller velocity, say 4.70 miles per second in the required direction at this range, it will enter an elliptical orbit having an apogee distance of 4,300 miles. The ratio $4.70/4.74$ is equal to $\sqrt{1-e}$ where e is the eccentricity of the orbit; e is therefore very nearly 0.017. The perigee distance will then be $(0.983 \times 4,300)/1.017$ or 4,158 miles. In these circumstances, therefore, an error of only about 1 per cent in orbital velocity will bring the satellite to within 160 miles of the earth's surface. On account of atmospheric drag, it would, of course, be unlikely to survive more than a fraction of a complete circuit of the earth. This example illustrates the exacting nature of the conditions for a successful launch.

The apogee distance r_a in an elliptical orbit is, from Figure 6 equal to $(1+e)/(1-e)$ times the perigee distance r_p ; the major axis of the orbit, which is equal to $(r_a + r_p)$, is therefore $2r_p/(1-e)$. The minor axis is $\sqrt{1-e^2}$ times the major axis, that is $2r_p \times \sqrt{1+e}/(1-e)$. The total area of the orbit is accordingly $\pi r_p^2 [\sqrt{1+e}/(1-e)]$. From the law of equal areas, it can now be proved that the periodic time is

$$T = 0.02 \times \sqrt{r_p^3} \times \frac{1}{\sqrt{1-e^3}}.$$

It follows that if we regard the horizontal scale of Figure 1 as the perigee distance, the value of T is given by multiplying the right-hand vertical scale by a factor $1/\sqrt{1-e^3}$.

Parabolic Orbits

The eccentricity of the orbit can be increased by making the displacing velocity greater: as e increases, apogee becomes more

and more remote. If a satellite having an orbital velocity v_0 , which suffices to keep it in a circular path of radius r_0 , is given an additional outward velocity v_d , it will, as has been shown, follow a new orbit having an eccentricity $e = v_d/v_0$. Its apogee distance will be given by $r_a = r_0(1 + e)$. It is clear, then, that if v_d is made equal to v_0 , the apogee distance will be infinitely great, so that the satellite will continue to recede from the earth. Its orbit can, in these circumstances, be regarded as having an eccentricity $e = 1$. The curve which it follows is a parabola.

We may use precisely the same method for determining points on a parabolic orbit as that employed in the problem considered earlier. Figure 7 shows an example of this when a satellite

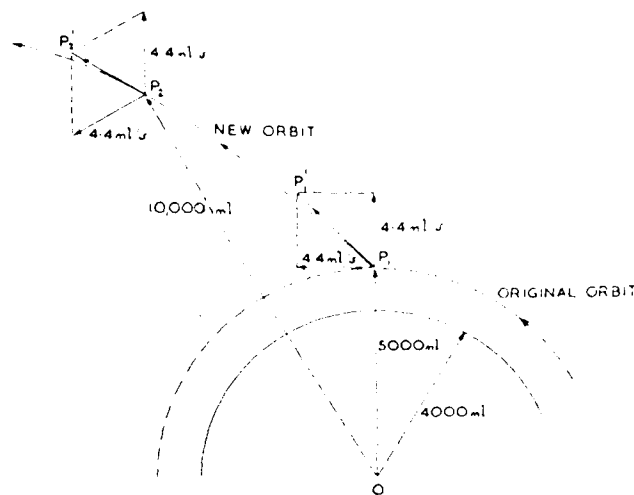


Fig. 7. *Parabolic orbit.* If the outward displacing velocity is equal to the original velocity in the circular orbit, the satellite enters a new orbit of unit eccentricity. This is the parabola P_1P_2 .

having an orbital velocity of 4.4 miles per second in a circular path is displaced therefrom by giving it an additional outward velocity of 4.4 miles per second.

It has been established that a satellite at a distance r miles, having a velocity v_p miles per second in a direction at right angles to the line from it to the earth's centre, will pursue an orbit that satisfies the equation

$$v_p = \frac{311}{\sqrt{r_p}} \times \sqrt{1 + e}.$$

In order that it should escape altogether from the vicinity of the earth, v_p must be large enough to make e unity. The condition for escape is therefore

$$v_p = \frac{311}{\sqrt{r_p}} \times \sqrt{2}$$

miles per second. To bring about the escape of a satellite which is travelling in a circular orbit, therefore, its velocity must be increased by a factor $\sqrt{2}$, or about 1.4, in the direction of original travel. Any smaller increase will cause it to follow a closed orbit in the form of an ellipse.

Conclusion

This article has described a very simple but somewhat unusual method for determining the orbit of a satellite. Moreover, it has been shown from the most elementary considerations how the orbit may be modified. This is a matter rarely treated in textbooks on mechanics: it has, of course, a special interest in relation to earth satellites. The reader who cares to experiment with the method will find no great difficulty in adapting it for the solution of a wide range of problems on planetary motion, even though his mathematical knowledge is limited to a little algebra and geometry. Indeed, the fascinations of armchair space navigation are available to any who may care to apply the simple principles that have been here outlined.

NUCLEAR COLLISIONS

P. E. HODGSON

PART II

IN the first part of this discussion of nuclear collisions (*Science News* 45) many of their characteristics were compared with those of a collision between two billiard balls. There is often quite a close analogy between what happens on a large scale and what happens in the nuclear domain, and an understanding of one can lead to an understanding of the other, provided that the limitations of the analogies are recognized.

But nuclear collisions are much more complicated and surprising to the physicist than are collisions between billiard balls. There are several aspects that have only the most tenuous relation to events on the billiard table. These are the phenomena of meson production, polarization, and nuclear disintegration.

MESON PRODUCTION

If billiard balls were made of flint instead of ivory, and if they collided violently together, we might see in a darkened room some sparks produced at the moment of impact. Something like this happens in nuclear collisions. When the energy is high enough, mesons are produced, and these, like the sparks, last only for a very short time.

But there is this difference: with billiard balls the sparks are little incandescent chips broken off the balls and heated by the violence of the impact, while in the nuclear case the mesons come entirely from the energy of the interacting particles, which cannot be chipped, however high the energy.

In Part I it was shown that the total cross-section, or probability of nuclear collision, varied with the energy, becoming less on the whole with increasing energy. At about 200 MeV

(million electron volts), however, it begins to level off, and remains constant until about 400 MeV, when it begins gradually to rise again. This is shown in Figure 1.

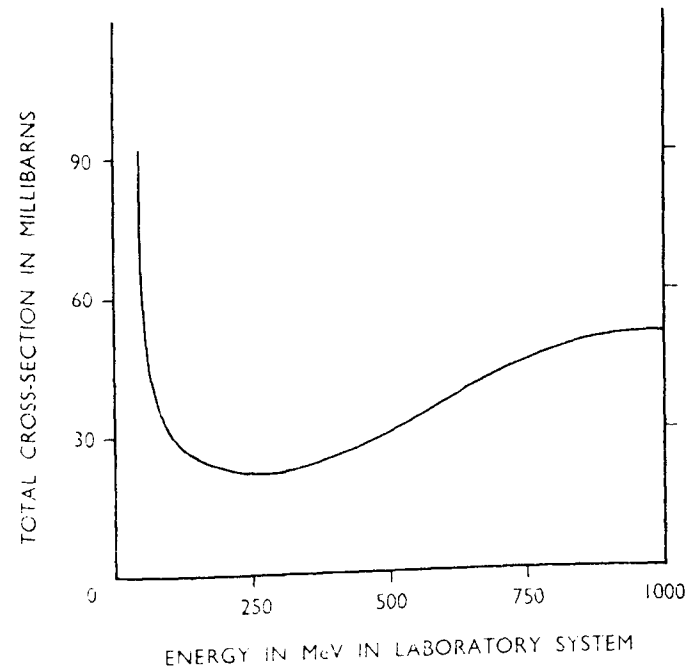


Fig. 1. The total cross-section for proton-proton collisions, showing the steady fall at lower energies and the subsequent rise owing to meson production.

At this point a digression on mesons is necessary before it is possible to understand this process of meson production. The story of how in 1932 they were postulated by Yukawa to account for the forces holding the nucleus together, and how they were discovered in the cosmic radiation by Anderson and Neddermeyer in 1936, has been told by Terry Price in *Science News* 16. Subsequently other types of mesons were found,

first in the cosmic radiation, and then among the products of the nuclear interactions produced by the particles accelerated to very high energies in cosmotrons. Their properties have been extensively studied and a summary of them will be found in the article on 'Fundamental Particles' in *Science News* 31. Although written only three years ago this summary is now seriously out of date, as many new discoveries and measurements, especially of the heavy mesons and hyperons, have been made since then.

Mesons can be divided into two classes, those that are produced in energetic nuclear interactions, and those that come only from the break-up of other mesons. As the energy of a nuclear interaction is increased, the first type of meson to be produced is the π -meson, or pion as it is often called. Pions can be either neutral, or positively or negatively charged. Neutral pions have a mass about 260 times, and charged pions about 274 times, that of the electron.

When a meson is produced, its mass must be taken from the energy of the collision, in accordance with Einstein's relation $E=mc^2$. Now the energy equivalent of an electron is just a little more than half a MeV, so about 130 to 140 MeV are required before a meson can be made. This energy must, however, be available in the centre-of-mass co-ordinate system of the colliding particles, and not just as energy in the co-ordinate system of the laboratory. Since for proton-proton collisions one energy is just about half the other, the energy of the incident proton must be at least 260 MeV before pions can be produced on energy grounds alone. In practice, the probability of pion production is appreciable only when the energy is substantially greater than this theoretical minimum. This accounts for the rise in the proton-proton cross-section in the region of 400 MeV owing to pion production. At higher energies two or more pions can be produced at once. At very high energies, whole showers of mesons are often produced. The higher the energy the narrower is the cone into which these mesons are emitted.

Neutral pions only live for a very short time before spontaneously transforming themselves into two energetic gamma

rays. The mean life of a neutral pion is about 10^{-14} or a hundred million millionth of a second. These gamma rays can, if their energy is high enough, produce pairs of electrons: these can easily be detected so the properties of neutral pions can be studied.

Charged pions tend to live rather longer; their mean life is about 10^{-8} or a thousand millionth of a second. They then transform themselves into a neutrino and another type of meson called a μ meson or muon. Muons hold the record for longevity in the meson family; they last, on the average, for about two millionths of a second before transforming into two neutrinos and an electron. Charged pions, muons, and electrons can all be detected quite easily by suitable instruments.

At still higher energies, heavy mesons, having the mass of about a thousand electrons, are also produced in nuclear collisions. Using arguments similar to those given above for pions, we find that this process may be expected to occur when the energy of the incident protons substantially exceeds 1,000 MeV; this has been confirmed experimentally. There are several different types of heavy mesons, each of which undergoes a characteristic series of transformations.

Quite recently an important discovery was made that had long been anticipated on theoretical grounds; this is the detection of the anti-proton. The mass of a proton is about 1,836 times that of an electron, so it might be expected that protons, both positively and negatively charged, would be produced in the collisions of protons of about 5,000 MeV. Positive or ordinary protons are well known as the nuclei of hydrogen atoms, but until recently there had been no experimental evidence for the existence of their negative counterpart. Theoreticians have, however, always preferred to have their theories symmetrical with respect to charge; that is, theories that would be unchanged if the signs of the charges were reversed throughout were preferred. This hypothesis had been verified for electrons and pions, so they had good reason to hope that it would hold also for protons, but until machines able to accelerate protons to about 5,000 MeV were completed it was not possible to test this.

When, however, the cosmotron at the University of California was completed and began to produce protons of 6,200 MeV, attempts were made as soon as possible to detect the antiproton, and they were successful in 1956.

POLARIZATION

We are all familiar with the way collisions between billiard balls are influenced by the spin we give to the cue ball; the same is true in nuclear collisions. In Part I it was shown that nuclear spins must be taken into account when interpreting these collisions by averaging over all possible spin states. This raises the question whether we can obtain beams of particles with their spins all pointing one way instead of having the usual beams with a mixture of spin states. If so it would be interesting to see what happens when the polarized beams collide with other nucleons; it would give us additional data with which to check our theories.

In 1948 several ways of doing this were suggested by Schwinger. It was shown that nuclear particles that have just collided, or have been emitted from a nuclear reaction, are very often polarized. But there is no way of finding out directly whether a beam of such particles is polarized; to do this a second scattering is necessary.

If an unpolarized beam is scattered, the same numbers of particles go in all directions making the same angle with the incident beam: we say the differential cross-section is azimuthally symmetrical. This is not so with polarized beams, and the amount of the polarization can be calculated from the numbers of particles going in different directions. These two possibilities are illustrated in Figure 2. If we simply measure the numbers of particles scattered to the right and to the left through the same angle θ , then, if the beam is unpolarized they will be equal, but unequal if it is polarized. The fractional inequality gives us a measure of the degree of polarization.

This is a difficult experiment to perform because two successive collisions are necessary, one to make the polarized beam and the other to show that it is polarized. At each collision most

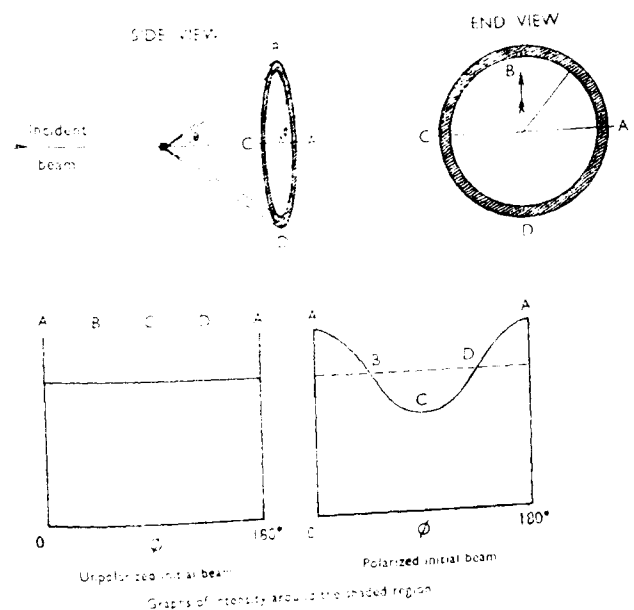
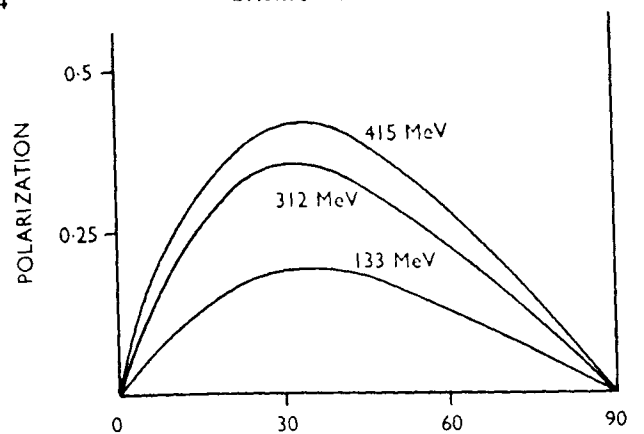


Fig. 2. Diagram showing the numbers of particles from polarized and unpolarized incident beams scattered into different azimuthal angles ϕ . Drawn for one scattering angle θ only. The arrow shows the direction of the polarization of the incident beam.

of the particles go in directions that are useless for the purposes of the experiment, so the initial beam has to be very intense if there is to be enough of it left after two collisions to allow accurate measurements to be made.

Improvements in cyclotron design have made possible the production of sufficiently intense beams, so in 1952 Oxley and his colleagues were able to demonstrate the production of a polarized beam for the first time. This has opened up a great new field of experiments in nuclear physics.

The results of some typical determinations of polarization in proton-proton collisions by a double scattering experiment are



SCATTERING ANGLE IN CENTRE-OF-MASS SYSTEM

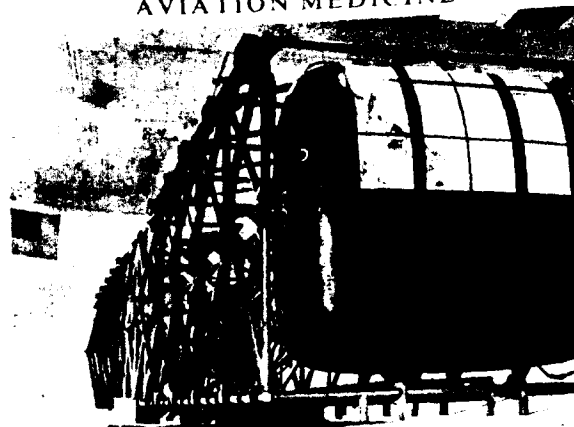
Fig. 3. Polarization in proton-proton collisions

shown in Figure 3. The polarization is plotted for several incident energies as a function of scattering angle in the centre-of-mass system.

To see what sorts of experiments can be done it is very useful to imagine a polarized beam as represented by a vector or arrow; the direction of the arrow shows the direction of polarization and its length the degree of polarization. This is correct only for beams of particles of spin one-half, but since protons and neutrons both come into this category, many of the most useful experiments can be easily visualized. (It is much more difficult to visualize the polarization of beams of particles of spin one, like deuterons. For them it is necessary to use a tensor as well as a vector to specify their polarization.)

We can conveniently think of an unpolarized beam as composed of equal numbers of particles with spin up and spin down, so that the resultant polarization is zero. Then, if such a beam hits a target that scatters more particles with spin 'up' to the left, and more particles with spin 'down' to the right, the emergent beams will be polarized as shown in Figure 4.

AVIATION MEDICINE



1. Looking along the sixty-foot arm of the man-carrying centrifuge at the R.A.F. Institute of Aviation Medicine at Farnborough. Accelerations of up to 30 G can be produced. (Air Ministry Photograph, Crown Copyright Reserved)



2. Subject seated in the cabin of the man-carrying centrifuge. A recording system picks up minute electrical impulses which are transmitted via the arm to the control room. These record changes in the subject's physiological condition. (Air Ministry Photograph, Crown Copyright Reserved)



3. The control room showing (centre) the apparatus on which the changes in the physiological condition of the subject in the rotating cabin are recorded. (*Air Ministry Photograph, Crown Copyright Reserved*)



4. Inside the Decompression Chamber at the R.A.F. Institute of Aviation Medicine. (*Air Ministry Photograph, Crown Copyright Reserved*)

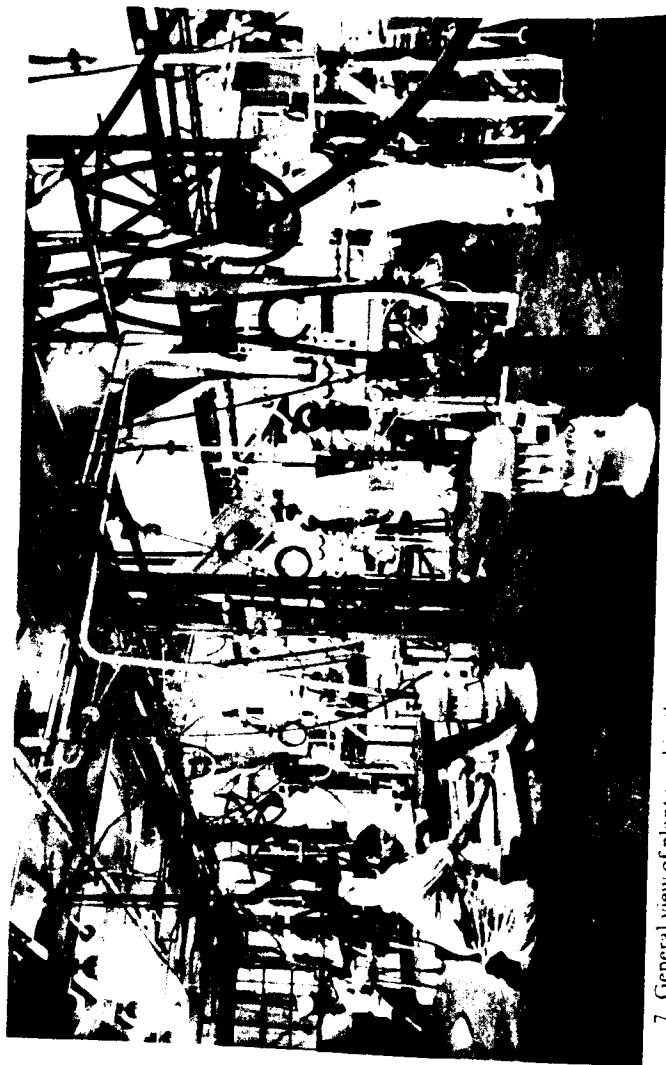


5. In the Climatic Research Laboratory at the R.A.F. Institute of Aviation Medicine. The Institute's work is chiefly concerned with the study of the physiological effects of flying under extreme atmospheric conditions in order to provide better protection for aircrews. Here a subject is adjusting an experimental oxygen mask. (*Air Ministry Photograph, Crown Copyright Reserved*)

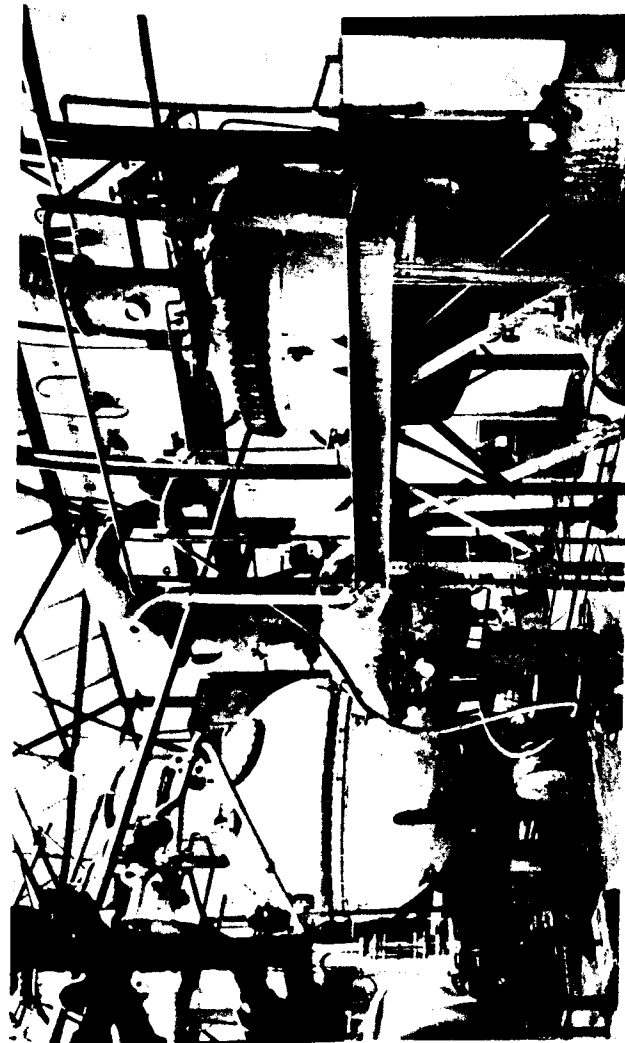


6. Testing flying clothing in a climatic research chamber at Farnborough. The subjects here are being tested at 120° F and 45 per cent humidity. (*Air Ministry Photograph, Crown Copyright Reserved*)

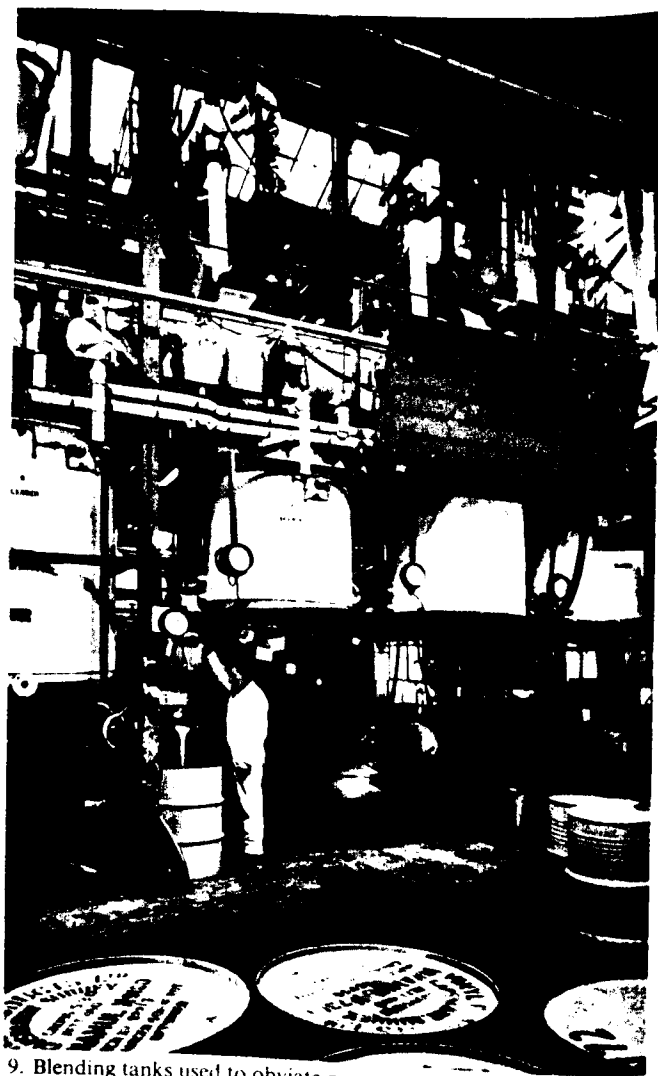
PLASTICS IN EMULSION



7. General view of plant used in the preparation of polymer emulsions. (Courtesy of Vinyl Products Ltd)



8. Distillation plant. Used for distillation of monomers from non-volatile inhibitors. (Courtesy of Vinyl Products Ltd)



9. Blending tanks used to obviate minor variations in the properties of successive batches of product. (Courtesy of Vinyl Products Ltd)



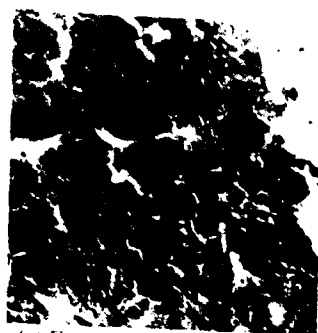
10. Polymerization reactors in operation, showing the upper part of the vessels with stirrer mechanism and condensers. (Courtesy of Vinyl Products Ltd)



(a) Touch dry



(b) One day after application

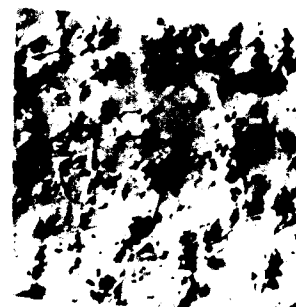


(c) Four days after application



(d) Eight days after application

11. Electron micrographs of film of plasticized polyvinyl acetate latex containing 20 per cent dibutyl phthalate: $\times 5000$, shadowing angle 15° . (Courtesy of the Research Association of the British Paint, Colour, and Varnish Manufacturers)



(a) One day after application



(b) Three days after application



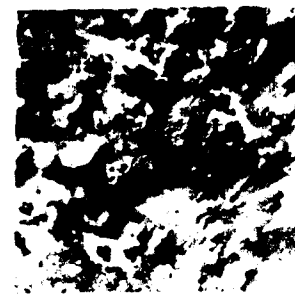
(c) Twenty-eight days after application



(d) Same as (c) but after washing



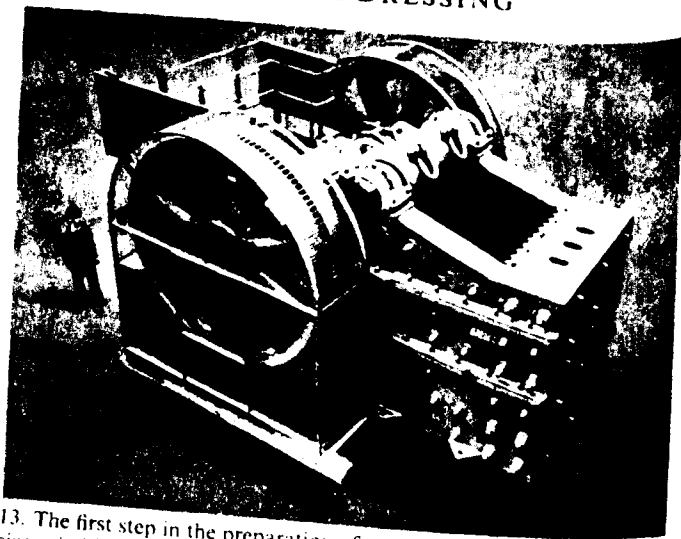
(e) Ten weeks after application



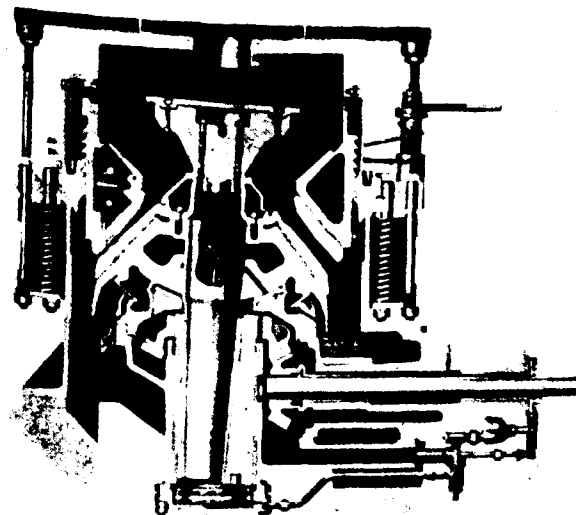
(f) Twenty weeks after application

12. Electron micrographs of surface replicas of polyvinyl acetate paint: $\times 5000$, shadowing angle 15° . (Courtesy of the Research Association of

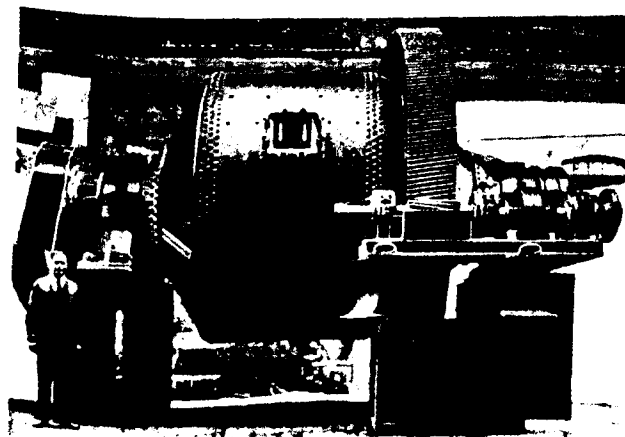
MINERAL DRESSING



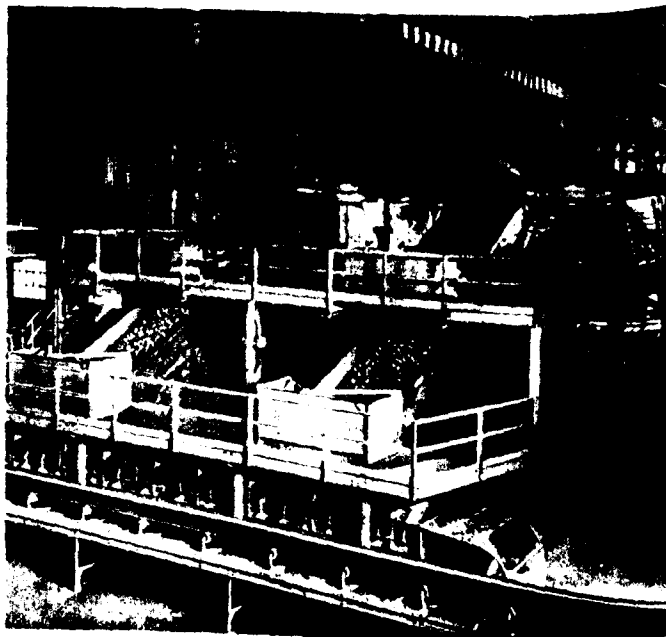
13. The first step in the preparation of ore is to reduce it to particles of size suitable for subsequent processing or concentration. This is usually carried out by dry crushing in a 'jaw crusher' of the type shown (Courtesy of the Fraser and Chalmers Engineering Works, Erith)



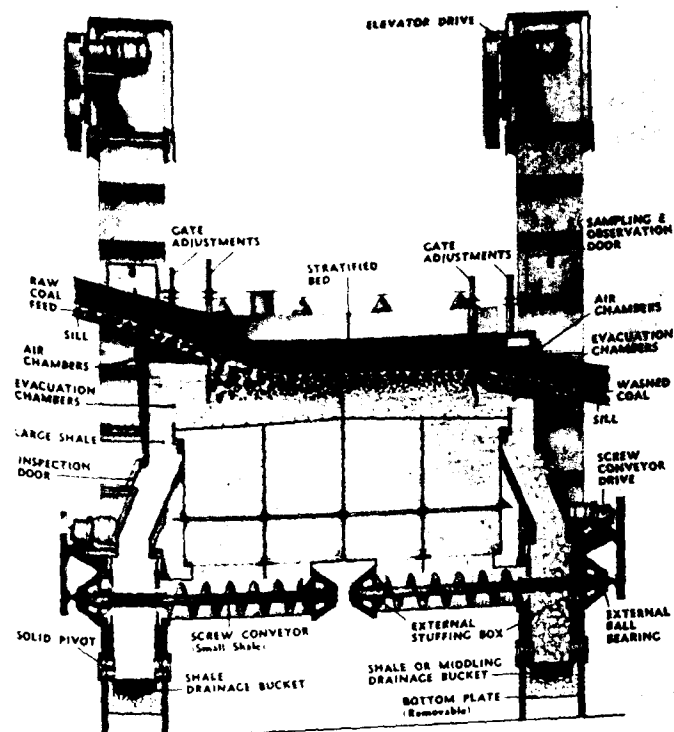
14. The crushed product from the primary breaker may be passed to a gyratory crusher. Shown here is the short-head Symons crusher. (Courtesy of the Fraser and Chalmers Engineering Works, Erith)



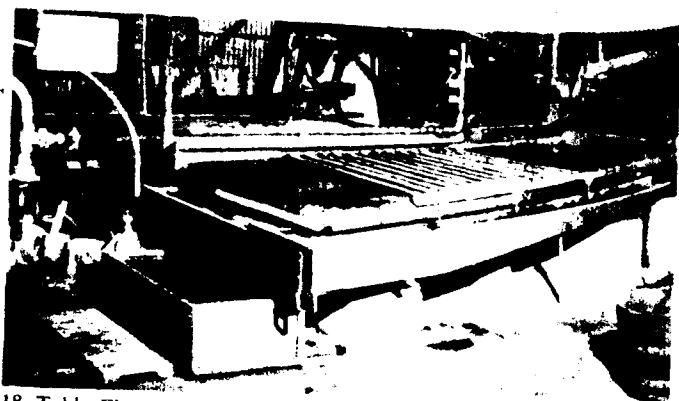
15. To release the valuable minerals and prepare run-of-mine ore for flotation or other methods of concentration, the dry crushed product is ground - nearly always under water. Nowadays, this is usually carried out in a tube or ball mill. (Courtesy of the Fraser and Chalmers Engineering Works, Erith)



16. In order to select processed ore of the required degree of fineness at the various stages of the process, the product may have to be screened. The plant illustrated is handling iron-ore. (Courtesy of the Fraser and Chalmers Engineering Works, Erith)



17. Baum Coal Jig: sectional illustration of four compartment wash-box. Raw coal is fed in over the inlet sill on the left. As it passes through the washbox compartments it is cleaned, and the washed coal is discharged over the outlet sill on the right. Heavy shale falls and is removed through bottom discharge gates. (Courtesy of the Automatic Coal Cleaning Co., Carlisle)



18. Table Flotation may be used to handle particles too large to be suitable for froth flotation. As with froth flotation, reagents are added to the ground concentrate. An insoluble oil is also added to effect the selective oiling of the particles to be separated. (Courtesy of Holman Bros. Camborne)



19. Cell flotation units of the type used in a lead-zinc mine in Wales. Air passes down the vertical shaft to a chamber where it is sheared into bubbles and intimately mixed with the mineral pulp. The mineral-bearing froth is removed by the froth paddle at the front of the tank. (Courtesy of Knapp and Bates Ltd, London)

Nuclear Collisions

65

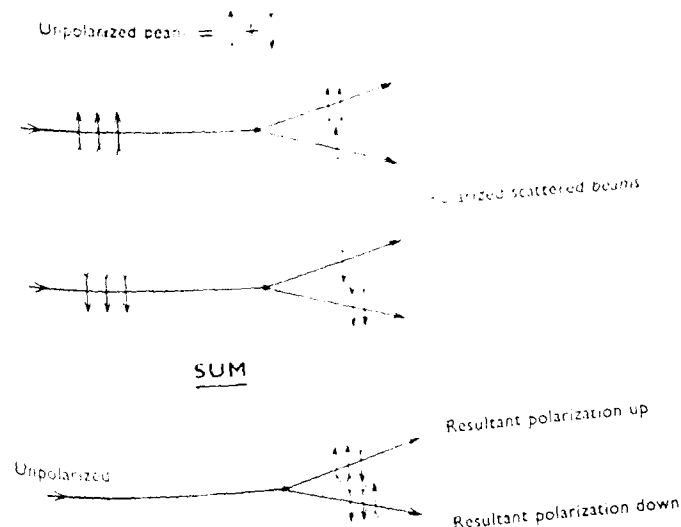


Fig. 4. Illustrating the production of polarized beams by a spin-sorting collision.

Here the target has acted as a spin-sorter. In other situations it can act as a 'spin-flipper' as well, and in the course of the collision flips the spins from one orientation to another.

It has been shown that the direction of polarization produced by the collision of an unpolarized beam is always perpendicular to the plane of the collision and to its direction of motion. This is called transverse polarization. But if the incident beam is itself polarized, the direction of polarization of the scattered

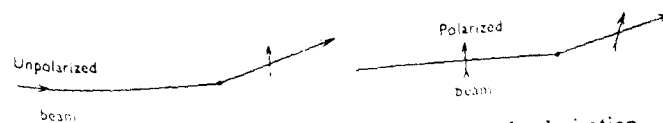


Fig. 5. Illustration of transverse and longitudinal polarization.

beam need not be perpendicular to its direction of motion; it will in general have a longitudinal component as well as a transverse one. Conversely, one scattering will only allow us to measure the transverse component of the initial polarization; to measure the longitudinal component a second scattering is necessary.

We can now look at some of the experiments that can be made with polarized beams. The simplest is the double scattering experiment that allows us to measure the polarization produced by scattering through different angles. The next in complexity is the triple scattering experiment. The first collision produces a polarized beam. The second (the one we are really interested in) alters this polarization, and the third measures the new polarization. Thus the first and third scatterers act as polarizer and analyzer. The two main types of triple scattering experiment, with the successive scattering planes either all in the same plane or at right angles to each other, enable quantities called the depolarization and the rotation of the polarization vector to be measured.

Similar measurements relating to the longitudinal polarization could also be made by quadruple and quintuple scattering arrangements. At present these experiments are prohibitively difficult because the available beams are not sufficiently intense, but fortunately the same information can be obtained in a much easier way. It has been found that if a polarized beam is passed through a magnetic field the direction of polarization is rotated with respect to its direction of motion. Thus by simply putting a magnetic field between the successive pairs of targets in the triple scattering experiments, measurements of longitudinal polarization can be made as well.

There is still another type of polarization experiment that is possible; we can ask what happens to the spin of the target nucleon. Suppose, for example, that a nucleon *A* collides with a nucleon *B*, and we knew that, after the collision, the spin of *A* pointed upwards. Whether this tells us anything about the spin of *B* is measured by the correlation coefficient. If we know for certain that *B*'s spin is up as well, then the correlation is one;

if certainly down, it is minus one; and if *B*'s spin is just as likely to be up as down it is zero. Intermediate values of the correlation coefficient, between one and minus one, are measures of probable knowledge. This is illustrated in Figure 6.

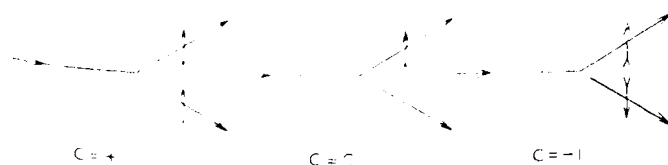


Fig. 6. The correlation coefficient.

The correlation coefficient can be measured by the arrangement shown in Figure 7.

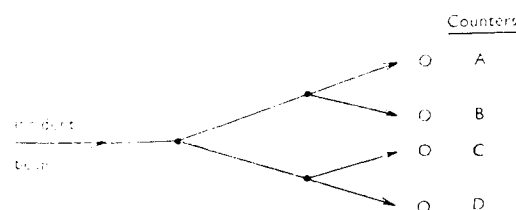


Fig. 7. Schematic diagram of the experimental arrangement for the determination of the correlation coefficients.

All the collisions take place in the same plane, and the four final beams are recorded by counters *A*, *B*, *C*, *D*. The pairs *AC*, *AD*, *BC*, *BD* are connected in coincidence; that is, a count is recorded for the pair *AC* only if a particle passes through *A* at the same time as another particle passes through *C*, and similarly for the other pairs. The correlation coefficient can be deduced from the four coincidence counting rates.

This experiment can be done either with a polarized or with an unpolarized incident beam, and successive collisions need not all be in the same plane.

The simple definition of the correlation coefficient given above must then be replaced by more exact quantum-mechanical reasoning, and when this is done it may be shown that there are five fundamentally different correlation experiments that can be made, two with an unpolarized incident beam and three with a polarized incident beam. At the present time some of these experiments are being tried in several laboratories, but so far no results have been published.

All the polarization experiments described so far have been done with unpolarized targets, or with targets composed of nuclei of zero spin, for which, of course, the concept of polarization has no meaning. To interpret these experiments it is, therefore, necessary to average over all the possible spin orientations of the target nuclei. It would clearly be an advantage if this could be avoided by the use of polarized targets. It is, however, extremely difficult to line up the nuclear spins in a solid; it can only be done by cooling to a very low temperature in an intense magnetic field. At present it seems that experiments with polarized targets will not be attempted, particularly as all we want to know can be obtained by the experiments already described.

The results of measurements of polarization phenomena can be calculated from any postulated nucleon-nucleon or nucleon-nucleus interaction in a way similar to that described for the differential cross-section in Part I. The greater variety of polarization phenomena provides a much more searching test of interaction, and many suggested interactions consistent with the cross-section data have had to be rejected because they are unable to account adequately for the results of polarization measurements.

NUCLEAR DISINTEGRATION

It is rare in billiards for one or other of the balls to be shattered, or even chipped, by the violence of a collision; but this happens quite frequently in nuclear collision if the energy is high enough. So far as we know the nucleons themselves are quite unbreakable, so this cannot happen in nucleon-nucleon collisions. All we can achieve by raising the energy of collision is the produc-

tion of more mesons and perhaps even more nucleons (proton-antiproton or neutron-antineutron pairs). But there is the possibility of break-up if a single nucleon collides with a cluster of nucleons, or if two nucleon clusters (or nuclei) collide. The only essential condition for such a break-up is that the energy of the collision is sufficient to separate one of the colliding nuclei into its parts or, in other words, the energy must be greater than its binding energy. These collisions are usually called inelastic, to distinguish them from elastic collisions in which the colliding particles just bounce off each other unchanged, except in energy and direction of motion.

The simplest inelastic nuclear collision is that of a nucleon with a deuteron. Since the binding energy of the deuteron, or the energy required to separate its constituent proton and neutron, is only 2 MeV, there is an appreciable probability of its being broken by collisions at quite low energies.

These collision processes are much more difficult to interpret than the nucleon-nucleon type because more particles are involved, but often they give information that could not be obtained in any other way. For instance, the collision of neutrons with deuterons can tell us something about neutron-neutron interaction; this could not be obtained directly because it is impossible to make a target out of neutrons alone since they are unstable. To a first approximation such a collision may be considered as the sum of a neutron-proton and a neutron-neutron collision, provided some allowance is made for the interaction between the two particles comprising the target deuteron.

In practice, if we find that the energy of a scattered nucleon is less than we would expect from Newton's conservation laws, we know that an inelastic collision has occurred; the missing energy has gone to break up one or other of the colliding particles. With deuterons the struck particle has only one way of breaking up, but in all collisions of more complicated nuclei there are several competing processes. For example, suppose a proton collides with a triton, a nucleus composed of one proton and two neutrons, then, quite apart from the possibility of meson production at high energies, six processes are possible.

Elastic collision: ${}_1\text{H}^1 + {}_1\text{H}^1 \longrightarrow {}_1\text{H}^1 + {}_1\text{H}^1$

Inelastic collisions: ${}_1\text{H}^1 + {}_1\text{H}^1 \longrightarrow {}_2\text{He}^4$

${}_1\text{H}^1 + {}_1\text{H}^1 \longrightarrow {}_2\text{He}^3 + {}_0\text{n}^1$

${}_1\text{H}^1 + {}_1\text{H}^1 \longrightarrow {}_1\text{H}^2 + {}_1\text{H}^1$

${}_1\text{H}^1 + {}_1\text{H}^1 \longrightarrow {}_1\text{H}^2 + {}_1\text{H}^1 + {}_0\text{n}^1$

${}_1\text{H}^1 + {}_1\text{H}^1 \longrightarrow {}_1\text{H}^1 + {}_1\text{H}^1 + {}_0\text{n}^1 + {}_0\text{n}^1$

Perhaps it is easier to visualize what is happening if we use a more pictorial representation, with a solid circle ● representing a proton and a hollow one ○ a neutron:

Elastic collision: ● + ○○ → ● + ○○

Bounce

Inelastic collisions: ● + ○○ → ○○

Stick-on

● + ○○ → ●○ + ○

Exchange

● + ○○ → ○ + ●○

Pick-up

● + ○○ → ○ + ● + ○

● + ○○ → ● + ● + ○ + ○

Break-up

The first of the inelastic collisions is the opposite of a break-up: the two colliding particles stick together and form one new particle (in this instance an alpha-particle). The second is a charge-exchange collision. The charge of the incident proton has jumped on to the triton or, in other words, the proton has knocked a neutron out of the triton and takes its place. In the third collision the incident proton has captured a neutron from the triton, thus forming two deuterons. The last two collisions are break-up processes.

All these processes compete with each other, though, since they require different amounts of energy, not all of them are possible at lower energies. In general, as the energy increases, the break-up processes become more probable, and the stick-on processes less probable.

These six possible types of collision are all that could conceivably take place, provided the combinations di-neutron ○○ and di-proton ●● are excluded because they are so unstable

but they are not found (if formed momentarily, they would break up at once).

The number of possible results for one of the simplest types of collision shows the vast number of nuclear reactions that can take place when heavier nuclei collide. This makes the analysis of such processes prohibitively difficult. Indeed, as soon as the number of reaction products is more than two, both experimental and theoretical difficulties become very great. Because of this, attention has so far been concentrated on the simplest examples of exchange and break-up collisions, and little headway has been made with the complicated ones.

It is very often true in particle physics that two extreme cases are amenable to experiment and calculation; the very simple cases with two or three particles and the case when there are so many particles that accurate results can be obtained even if the finer details of the process are neglected or averaged out. The second applies, for example, to the problems studied in the kinetic theory of gases. Even though we have only probable knowledge of the motion of any individual particle we can average out our ignorance and make quite accurate calculations of the pressure, volume, and viscosity of a given amount of gas. We can also calculate quite well what happens when two gas atoms collide. But if there are even three or four, the mathematics becomes hopelessly complicated, and problems can only be solved if circumstances allow plausible approximations to be made. Thus it is fairly easy to deal with very small or with very large numbers of particles; it is the intermediate situations that are most difficult.

The same is true in nuclear physics. So if we go from the simple collisions described above to collisions involving a large number of nucleons we once again find that simple concepts can be developed that enable the processes to be understood.

Such is the case when a very energetic nucleon collides with a fairly large nucleus containing some tens of nucleons. What happens has been described in some detail in an article on 'High Energy Nuclear Disintegrations' in *Science News* 33, so it will only be briefly summarized here.

The process can be divided into two stages, the cascade stage and the evaporation stage. The incoming nucleon enters the nucleus and collides with one or more nucleons in it, knocking them on and perhaps producing mesons if the energy is high enough. These secondary nucleons and mesons often make further collisions within the nucleus, until they finally emerge from it. After the passage of this meson-nucleon cascade, the residual nucleus is left with a large amount of energy. This energy is soon shared among all the nucleons in it. Now and then, however, a particle near the surface of the nucleus will receive enough energy for it to escape, and this process will continue until the excess energy of the nucleus is so low that no more nucleons can escape. The small amount of energy still there is usually dissipated by the emission of electrons and gamma rays.

Both stages can be fairly well treated theoretically. The collisions in the cascade take place at such a high energy that by comparison the forces that bind the target nucleons to the rest of the nucleus are negligible. The results previously obtained for free two-nucleon collisions may therefore be applied with little loss of accuracy. The evaporation stage can be treated statistically using one or other of the nuclear models described by Satchler in *Science News* 43. One of these is particularly simple to visualize: it compares nuclear disintegration to the evaporation of a drop of hot water. In this way the relative proportions of the different types of emitted particles, as well as their energy and angular distributions, can be calculated for comparison with experiment.

Of course many of these particles are polarized, and it would be interesting to study how the polarization of the reaction products depends on the energy and polarization of the incident particle. But this could be done only in a rather general way because the polarization of any particular particle cannot be measured; as shown in the previous section, it is only possible to measure the average polarization of a beam of particles.

CONCLUSION

By making use of analogies, in particular that between the collisions of billiard balls and those of nucleons and nuclei, a few of the most important features of nuclear collisions have been described. More detailed description and understanding is not possible without recourse to mathematical analysis. But it has been possible to get quite a long way without it, and to show what an extensive field of experimental and theoretical research is being opened up by the study of nuclear collisions.

AVIATION MEDICINE

G. W. G. SHARP

THE rapid advance in aircraft performance during the last few years has brought many problems for the physiologist to solve. These are caused by the stresses and abnormal conditions under which pilots and crews now have to fly. In this article we shall consider a few of these problems, and indicate the ever-increasing scope of what is commonly termed 'aviation medicine'.

Physiological studies were first applied to aviation during the first world war, when certain tests were devised to assess the fitness of an individual for flying duties. In point of fact the tests were not of great value because they tended to be tests of bodily fitness rather than tests of aptitude for flying. Nevertheless they did assess such qualities as reaction time, vision, and general alertness. These qualities, however, although important, bear no direct relationship to problems peculiar to flying – for example, airsickness. After 1918, very little further work was carried out until the middle thirties, when, with the threat of another war and the importance of air supremacy, medical research on the topic restarted. This time, however, research was applied more directly, and such factors as the effects of rarefied atmospheres, centrifugal forces, noise, cold, and glare were investigated. These studies proved their value when the war came, bringing with it an even more vital need for research, and the efforts of doctors and physiologists in the intervening years, and during the war, helped in no small way towards the attainment of eventual air supremacy by the Allies. This was particularly obvious in the days of the Battle of Britain, when the need to fly and manoeuvre at speeds and altitudes greater than the enemy introduced the two major problems of 'blacking out' and oxygen-lack. Before the war most civil and military flying had been carried out at altitudes where the concentration of oxygen in the atmosphere was compatible with life.

PROBLEMS ASSOCIATED WITH G*

At the present time, with aircraft regularly exceeding the speed of sound, it is clear that efficient anti-G measures are essential. It is well known that blacking out is due to pooling of the blood in the legs and lower abdomen, with a consequent blood shortage in the brain. This is brought about by centrifugal forces

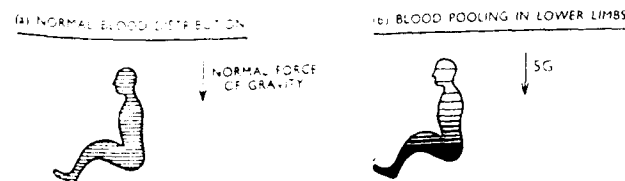


Fig. 1. *Black out due to high G.* Diagram to show the pooling of blood in the lower limbs due to the application of a high gravitational force in the direction head to feet. The resultant blood shortage in the head may suffice to cause loss of consciousness.

applied in the direction head to body, and occurs during sharp turns at high speed, or when pulling out of a fast dive. Normally the heart maintains a blood pressure sufficient to supply the brain with blood under most conditions, but it is not capable of responding quickly to centrifugal forces in the region of 4–10 G. As the force is applied, 'grey out', 'black out', and finally unconsciousness occurs, simply because the heart cannot pump blood to the brain against this force. The pilot feels 'heavy' and may have difficulty in keeping his head erect. The skin of the face is forced downwards, particularly on the cheeks and eyelids, and the lower jaw sags, the arms are difficult to move, while the lower limbs are congested, rigid, and cramped. Expiration must be forcible because the diaphragm is pulled down and the abdominal wall pushed out by the weight of the liver, stomach, intestine, and abdominal contents. In the relaxed subject, visual disturbances appear at 3–4 G, dimming and greying first, then

* G = terrestrial gravitational force taken as a unit.

black-out at 4-5 G, and unconsciousness at 5-6 G. As G decreases, consciousness returns with mental confusion and slight amnesia; vision returns after a few seconds, and commonly the only after effect is a feeling of tiredness.

The circulatory symptoms are due to the forced expansion of blood vessels below the heart, the failure of venous return, and the reduced pressure above the heart. At first research into problems associated with G was hindered by the lack of suitable recording apparatus and, in particular, apparatus for recording circulatory changes. Then, again, the reports given by flyers varied widely, and as high G causes mental confusion they could not be relied upon for accurate scientific work. Therefore the development of the oximeter was a great help towards further progress because it is capable of continuously recording the blood supply to the head.

This instrument consists of a lamp and a photoelectric cell, which when applied to the ear measures its opacity. The apparatus is attached so that the lamp is separated from the photoelectric cell by the ear itself; thus the amount of light reaching the cell depends upon the ear opacity. This in turn depends upon two factors, one constant, the bloodless tissue of the ear which accounts for some 70 per cent of the opacity, and the other, the relative opacity of the blood within the blood vessels. This latter is obviously variable, and a reduction in blood supply, or in the oxygen content of the blood, will cause a reduction in the opacity of the ear. A reduction of opacity will therefore signify a lowered supply of oxygen to the ear, and also, of course, to the head and brain. The oximeter can be calibrated and an accurate continuously recorded picture of the blood supply to the head easily obtained. With the improvement of the apparatus, it became possible to record simultaneously the blood supply to the brain, the heart rate, the electrocardiogram, pulse waves, arterial blood pressure at various levels, the response to light, and the gravitational force, so that progress became more rapid. It was found for instance that the effect of G is considerably influenced by the speed of onset, its duration, and the compensatory mechanisms of the body. Suppose for example that a per-

son is exposed to 3G for half a minute. The effects are quickly observed and consist of tachycardia (acceleration of the heart), decreased ear opacity, decreased blood pressure at head level, failure of peripheral vision, and probably retention of central vision. After some seconds, however, the compensatory mechanisms come into play. The reduced blood pressure in the carotid sinus brings about vasoconstriction and a raised blood pressure, with consequent return of peripheral vision due to the increased blood supply to the head.

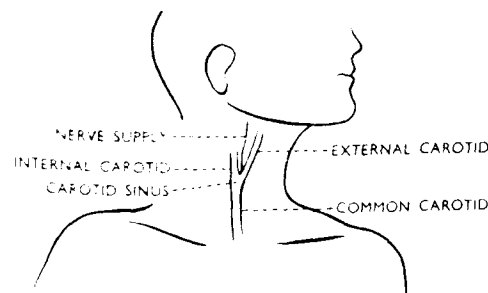


Fig. 2. *The Carotid Sinus.*

When experiments are carried out over shorter periods these compensatory effects are not seen, although on the other hand, if the exposure to G is too short, the full effect of the high G may not manifest itself. The continuous recording gear will therefore show the effects of varying gravitational forces applied for varying periods of time; any positional changes and protective devices can be tested physiologically for advantageous effects. The only remaining difficulty hampering rapid progress was the application of the gravitational force. The work could be carried out in aircraft, but the attainment and maintenance of high G values for long periods is impossible. This resulted in the development of a centrifuge designed to carry a human subject.

The reverse of blacking out, when force is applied in the direction body to head, is known as 'redding out'. It occurs if a

pilot attempts an outside loop. In this case the blood pools in the skull under pressure, vision is obscured, and prolonged high G values under these conditions result in eye and brain haemorrhage of a serious nature. This effect is, however, only met with rarely, and need not be considered further.

The amount of G required to produce blacking out varies considerably from one person to another, and pilots soon found that tensing muscles, grunting, shouting, 'bearing down' and similar procedures had a favourable effect in pull-outs and tight turns. All these favour the return of venous blood to the heart, by compression of arterioles with consequent increase of blood pressure, and by compression of veins and increased venous flow. The intra-abdominal pressure is raised, and large veins compressed; it is noticeable that a full stomach increases one's resistance to raised G.

In the absence of a 'human centrifuge', aptitude tests were carried out in the air. It was quickly realized that the pilot's posture had a considerable influence on his ability to resist centrifugal forces and consequent blacking out, this being lowest in the upright position and highest in the prone.

The prone position was clearly impracticable at the time (and, generally speaking, still is today) so a compromise had to be evolved. This consisted essentially of raising the rudder pedals

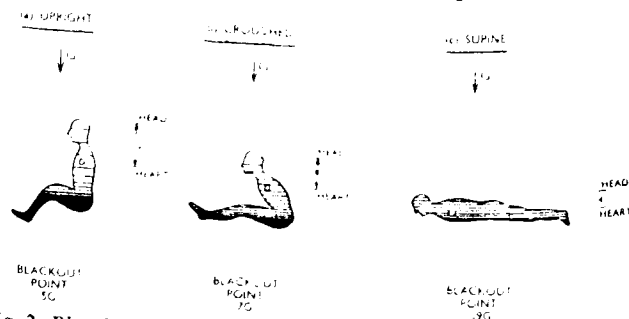


Fig. 3. Blood distribution due to the application of high G. Decreasing the effective heart to head distance increases the resistance to high G forces.

so that the feet were almost horizontal, and of making the pilot lean forward at an angle. The former reduced the pooling of blood in the legs, while the latter reduced the difference in height between the heart and the brain, thus lessening the force required to pump blood to the brain. It can be seen that in the prone position the heart and brain are on the same level, so that the heart does not have to work in direct opposition to the centrifugal force. This simple postural change was sufficient to raise the blacking out point from approximately 5G to the region of 7G. Further increases in resistance to blacking out were produced by pressure suiting. It was argued that if pressure was applied to the abdomen, then under conditions of high G pooling of the blood would be reduced. This led to the use of abdominal belts, which, however, were found to exert only a slight effect on G resistance; but they did minimize the downward movement and sagging of the internal organs.

Pressure on the legs alone gave very little protection, though when it was applied together with pressure on the abdomen the efficiency of the latter measure was almost doubled. Pressure suits were much superior to abdominal belts, and by gradual development and improvement they are now capable of raising the blacking out threshold by another 2G. They act by diminishing the peripheral arteriolar bed, increasing the peripheral resistance and thus increasing both the arterial blood pressure and the return of blood to the heart. The development of these suits and indeed of all anti-G work was speeded up considerably by the introduction of the human centrifuge. This consists essentially of two 'arms' fifty feet long, which can be made to revolve at any required speed. At the end of one arm is the test cockpit containing the subject, together with instruments which record both the stresses to which the subject is exposed and his reactions to those stresses. The cockpit may also be in the form of a pressure chamber, or a temperature chamber, as required. This type of centrifuge has greatly facilitated precise research into the problems of G under varying conditions of pressure, oxygen supply, and temperature which were not easily performed prior to its introduction.

For instance the application of prolonged G in a centrifuge has shown the two phases of effect quite distinctly; initially a fall in blood pressure, tachycardia, and failure of vision, followed by adaptation due to vaso-motor changes. The return to normal G is accompanied by a rise in blood pressure and blood supply to the brain, to levels above the normal, due to an overshoot effect of the compensatory mechanism. This results in some important effects. Exposure to a high G will protect to a certain extent against a second exposure some seconds afterwards owing to the activity of compensating mechanisms. On the other hand exposure to negative G (redding out) followed by positive G produces more severe reactions than the positive G alone, because the compensatory mechanism is stimulated to full effect on the blood pressure by negative G and this then aids the normal high G effect. The result is a profound fall in blood pressure and blood supply to the head. Tests in animals have shown that if the nerve supply to the carotid sinus is cut, the effects of high G are increased in severity, owing to the failure of the compensatory mechanism.

THE SUPPLY OF OXYGEN

High altitude flying introduces problems of oxygen supply and pressure conditions that are complicated by the possibility of emergencies. The effects of oxygen-lack due to increase of altitude are of course variable, depending upon the duration of the lack, upon the individual subject, and also upon the rate of increase in altitude. The latter, which is an illustration of acclimatization, is obvious when one compares the performance of mountaineers, who on Mount Everest have reached a height of over 28,000 feet without oxygen equipment, and a pilot, whose oxygen equipment fails at the same height, who will lose consciousness in two to three minutes. A pilot flying without oxygen will probably be affected by oxygen lack at about 6,000 feet, when an increase in the rate and depth of breathing will occur. The serious effects will not be seen until an altitude of about 12-15,000 feet is reached, when typical symptoms of anoxia will occur. They start with an impairment of physical

and mental processes associated with a dangerous feeling of well-being and confidence. At high altitudes circulatory changes attempt to compensate for the lack with an increased cardiac output and rate, all the senses are dulled, and finally at varying heights between 20,000 and 30,000 feet the sudden onset of unconsciousness is the prelude to paralysis and death, unless the oxygen lack is remedied.

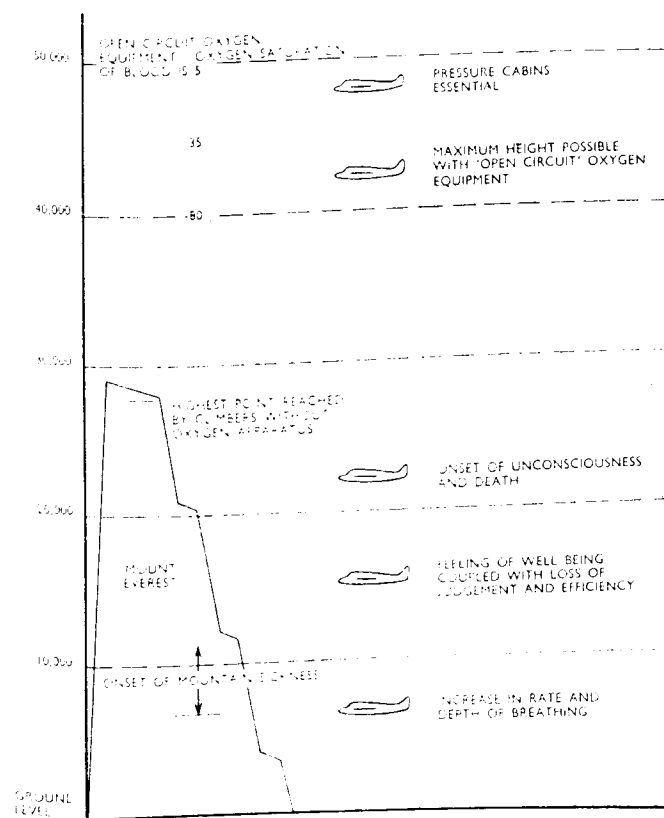


Fig. 4. Diagram to illustrate the effects of increasing altitude.

Clearly then, oxygen must be supplied to aircrews if they are to fly for any length of time at heights of over 10,000 feet. Open circuit equipment will suffice to give adequate blood saturation up to 30,000 feet, above which the partial pressure is too low for even pure oxygen to provide 'complete' saturation. The maximum height at which the open circuit equipment can be used is less than 45,000 feet. The following figures show this: breathing pure oxygen at 40,000 feet, the oxygen saturation of arterial blood is approximately 80 per cent. At 45,000 feet the saturation is 35 per cent, and at 50,000 feet the saturation has dropped to below 5 per cent.

The open circuit mask is clearly very wasteful of oxygen, and to provide oxygen for the crew of a bomber, say, would require extra equipment taking up space and increasing the weight, so that the bomb-load, range, or maximum altitude would suffer accordingly. Therefore a closed circuit was required that would conserve oxygen. This was brought about by passing the exhaled air through a soda-lime cylinder to remove carbon dioxide, and then remixing it with the incoming air supply. One great difficulty with this type of equipment is to prevent water vapour in the expired air freezing and blocking the supply lines. Obviously the ideal is a system controlling the oxygen flow according to the pilot's needs. This was provided by a mask containing a 'diluter' valve controlled by a barometer. The greater the altitude, the more oxygen supplied, and so on. However the lack of oxygen is not the sole disadvantage attached to high altitude flying.

LOW PRESSURE EFFECTS

The reduced pressure, apart from causing oxygen lack, will also cause distension of the intestines and other organs, owing to expansion of gases. This may in some cases cause severe pain. A further point that in recent years has changed from academic interest to practical importance is that, at an altitude of over 63,000 feet, the atmospheric pressure is so low that at body temperature blood and body fluids boil. These difficulties of high flying have been overcome to a large extent by the develop-

ment of pressurized cabins. Aircraft carry air compressors which keep the atmosphere in the cabins at pressures well above those of actual flight levels, so that no masks are required, there are no tubes to restrict one's movement, and people can talk quite freely. This is quite satisfactory for an airliner which cruises at a set altitude, but is impracticable for a military aircraft which may have to make rapid and considerable changes in altitude, and is designed to operate at heights greatly exceeding those of an airliner. The reasons for this are that a fighter aircraft, for instance, cannot be safely pressurized for a constant level, the weight of compressor machinery reduces performance, and a practical limit for pressurization seems to be close to 80,000 feet. Above this limit, the sealed cabin with its own 'built in' atmosphere is necessary.

The equipment of a modern fighter pilot serves to illustrate most of the problems associated with low pressure and oxygen-lack. For instance, the gradual reduction in pressure from ground level to about 30,000 feet is compensated for by an automatic valve mechanism which supplies an air and oxygen mixture to the pilot. The oxygen content automatically increases as the altitude increases, until eventually 100 per cent oxygen is being supplied. The cabins are pressurized to a certain extent, but at a flight level of 55,000 feet the cockpit pressure may only be equivalent to 45,000 feet. Now, as stated previously, the oxygen saturation of arterial blood (when breathing pure oxygen) is only 35 per cent at 45,000 feet and this is insufficient to support bodily functions. The oxygen has, therefore, to be fed to the lungs under pressure. Under normal circumstances a person creates a low pressure in the lungs by contraction of the diaphragm and expansion of the chest walls, causing air to be drawn in, while the expiratory process is purely a passive movement, i.e., the relaxation of the diaphragm, fall of the ribs, and recoil of the lung passages, all combine to force air out of the lungs. With pressure breathing, however, the normal process of respiration is reversed. Inspiration is passive, owing to the oxygen under pressure forcing its way into the lungs, while expiration is active, and the lungs have to be forcibly cleared of

oxygen and carbon dioxide against the oxygen pressure. This is not at all easy to accomplish at first; it makes speech difficult and flyers find it tiring.

Hyper-ventilation has been found to occur in pilots, particularly under conditions of fear or stress. This excessive breathing disturbs the O_2/CO_2 balance, and increases the alkalinity of the blood. Dizziness, loss of co-ordination, and muscle spasm, particularly of the hands, may result, with obvious consequences to the pilot of a high-speed aircraft. So far the only solution to this problem, which is believed to have caused a number of accidents, is to make the pilot aware of the symptoms so that he can either control his breathing consciously or temporarily hold his breath.

At heights above 50,000 feet, loss of pressure in the cockpit, which might occur owing to enemy gunfire or simple compressor failure, constitutes a serious threat to survival, because consciousness is lost in less than fifteen seconds. At one time it was thought that explosive decompression of this nature would cause serious lung injury owing to distension and tearing, but by experiments upon animals and man this has been shown to be untrue – providing that the glottis remains open. Pressure breathing has been shown to be impossible at these heights, because the pressure in the lungs, being far greater than the external pressure, would place too much strain on the lungs and chest. The first means of combating this emergency was the use of a pressure waistcoat which inflated automatically when the cockpit pressure was lost. The waistcoat provided the necessary external pressure to the chest and enabled pressure breathing to maintain the intra-pulmonary gas pressure and arterial oxygen saturation at adequate levels. A logical development was a complete skin-tight suit, fitted with inflatable tubes, but these suits proved to be so rigid and uncomfortable that they could only be worn for short periods. They did, however, allow pilots time to descend to a safe altitude after sudden decompression or loss of atmosphere. Recently, true pressure suits have been designed so that the air under pressure is next to the skin, and the pilot in fact is enclosed in a miniature pressure cabin. Naturally the

apparatus necessary to pressurize a suit with its small volume is much lighter and more practicable than that required for a pressurized cabin, and is thus more suitable for use in light fighter aircraft.

Emergencies in stratosphere flying, which entail escape from an aircraft at extreme heights, have again brought a considerable number of physiological difficulties. As the aircraft may be out of control and flying at high speed, the pilot must be ejected from the cockpit in order to avoid hitting the tail of the aircraft. He is ejected by an explosive charge and on leaving the aircraft has to face the blast from the slipstream, lack of oxygen, and intense cold due to the rarefied atmosphere and high altitude, and a low external pressure at which his body fluids might vaporize. In addition, again owing to the rarefied atmosphere, the speed of descent may be sufficiently great to preclude the safe opening of a parachute. By experimenting it has been shown that decelerating forces of up to 30 G may be produced owing to a parachute opening soon after bailing out at 35–40,000 feet. This means that the pilot must wait until he slows up in the denser air of lower altitudes before he opens his parachute. In this free fall period, dangerous centrifugal forces may build up as a result of the pilot and ejection seat spinning at high speed. Since 1944, when the injuries received by the first pilots trying to escape from jet aircraft were carefully analysed, all these problems have been investigated by doctors and physiologists, some of whom have carried out large numbers of test jumps. It was obvious that the body could not withstand the blast from the slipstream, and even with the advent of the ejection seat, the wind was still sufficient to tear away helmets and oxygen masks essential for the free fall period. Facial injuries were also serious, and eventually a canopy was invented which is pulled down over the head until it reaches the level of the sternum. At this point the canopy triggered off the explosive charge, and the force of ejection ensured that the canopy in turn pulled the fighter over the head and face. This provided adequate protection for the initial stage of escape. Further work developed suitable pressure suits and oxygen sets for the drop, while clothing

sufficient to insulate the wearer from the cold was incorporated with the suits. It would appear that an automatic parachute is desirable, set to open at say 10,000 feet, in case the airman is injured or unconscious, and also in case he has blocked out owing to spinning. This latter problem has several possible solutions, including ejection seats fitted with short wings which open after ejection and stabilize the fall, and ejection of the cockpit as a whole from the plane, when, owing to the approximately cylindrical shape, less spinning is likely. This may seem to be more nearly the work of aircraft engineers and designers than physiologists, but in fact the latter are organizing and planning this development of ejection apparatus and in many cases have acted as their own 'guinea pigs'. Indeed, many factors that we accept today without much thought are due to physiological testing and development in close co-operation with instrument makers, clothing manufacturers, aircraft fitters, designers, physicists and so on. Among these are such things as the layout of instruments for highest efficiency and for lessening fatigue, lighting arrangements, seat design and positioning, night vision adaptation aids, suits that allow the passage of water vapour (perspiration) but which remain waterproof in sea-water, the design of life jackets that keep the nose and mouth of an unconscious individual above water, and so on.

STRATOSPHERE HAZARDS

Finally, with an altitude record set at over 90,000 feet, an altitude at which more 97 per cent of the mass of the earth's atmosphere is below, and with the development of rocket propulsion, which will enable man to leave the earth's atmosphere, it may be appropriate to consider very briefly two of the problems that flight on the fringe of, and outside, the earth's atmosphere will produce.

Cosmic rays are clearly important. It is well known that the earth is constantly bombarded by atomic nuclei, which are rendered harmless by the atmosphere. Primary rays may penetrate to an altitude of 70,000 feet and the number of penetrating rays increases up to 120,000 feet, an altitude to which

most nuclei penetrate. It is believed that most of the penetration is harmless, but that the heavy particles which comprise some one per cent of the whole may be dangerous. At the moment the danger is slight because of the small number of people involved, and because of the short duration of the really high altitude flights, but the need for research is of paramount importance for the future. It has been shown that the life expectancy of animals is reduced by exposure to densely ionizing rays, while possible genetic dangers are realized by many people.

Assuming that man can be shielded from cosmic rays, and that he flies outside the earth's atmosphere, the condition of 'weightlessness' will arise. This will occur when no gravitational forces are acting, as for example when a projectile floats in space without motor propulsion. Air-breathing engines are not capable of operating at altitudes much greater than 90,000 feet, but now that rocket motors are being used in aircraft, higher altitudes and longer periods of weightlessness may be experienced. This phenomenon has been investigated to a certain extent, by flying in an arc of known curvature and at calculated speeds so that the centrifugal force equals and cancels out the force of gravity. So far weightlessness has only been produced for periods of up to thirty seconds, and various effects have been reported. These are confusion, and sensation of spinning, and naturally a lack of co-ordination. The body is adapted, of course, for the normal force of gravity. Will people under conditions of weightlessness feel as though they are constantly falling, or will they adapt themselves? Again, the tests undertaken so far have produced weightlessness for short periods only. Will long exposure result in adaptation or further complications such as malfunction of the heart and disorder of the circulatory system? The answer to these and many other problems of aviation medicine will only be found by much further research.

RESINS IN EMULSION

H. WARSON

MOST people are familiar with rubber latex, the milky liquid obtained from *Hevea Brasiliensis*. This is a natural form of rubber in which the rubber is suspended as very fine particles in an aqueous medium. About thirty years ago the idea of making a synthetic resin in the form of a latex was developed on an industrial scale. The term 'latex' as used here is almost synonymous with 'emulsion', the term most frequently used in Britain; an emulsion may be defined as a dispersion of drops of one liquid in another, so made that the system remains stable for weeks rather than minutes.

Emulsions are of two main types, oil in water, and water in oil, but only the former need concern us here. For an emulsion to remain stable and the two phases not to separate out, some form of emulsifying agent must usually be present. This may be one of several types, but sodium stearate $C_{17}H_{35}COONa$, a constituent of ordinary soap, may be taken as typical. Its molecule consists of a hydrophobic part, a long hydrocarbon 'tail', and a hydrophilic part, a sodium ion.

Emulsifiers function by their surface activity at the interface of the water and of the non-aqueous droplets, the hydrophobic part being attracted towards the latter, and the polar hydrophilic group towards the water. Hence the term 'surface active agents' or 'surfactants' which is often used to describe these emulsifiers.

Soaps are probably the oldest emulsifying agents, but now a very large number of synthetic agents have been prepared. Sulphate and sulphonate groups very often form part of the molecule of a synthetic emulsifying agent, as in sodium lauryl sulphate or the sodium salt of dodecylbenzene sulphonate acid, and the hydrocarbon group may be in the 'positive' instead of the

'negative' part of the molecule as in cetyl ammonium bromide.

There is a further group described as non-ionic emulsifiers. They rely on hydroxyl or ethoxy groups to achieve the balance due to the hydrophilic part of the molecule. Glycerol monostearate and polyethylene oxide lauryl ether are examples.

Emulsions can also be stabilized with some water soluble materials of high molecular weight, e.g., dextrin, gum arabic, gum tragacanth, and certain synthetic materials such as methyl cellulose and polyvinyl alcohol.

When they can be prepared, mixtures of emulsifying agents are often much more effective than their individual components; a very effective system being one in which an emulsifying agent is combined with a colloid. The former breaks down the oil phase into small particles of the order of about 1 micron (1μ) or less, and the colloid imparts stability.

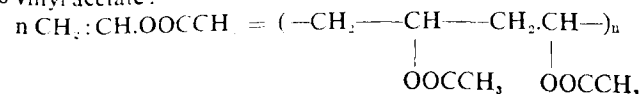
VINYL RESINS

Vinyl resins are formed by polymerization of unsaturated monomers. The simplest monomer is ethylene, $CH_2:CH_2$. Other well-known monomers include vinyl chloride, $CH_2:CHCl$, methyl methacrylate, $CH_2:C(CH_3)COOCH_3$, and vinyl acetate, $CH_2:CHCOOCH_3$, from which a resin is formed that enjoys great popularity as an adhesive and as a basis for emulsion paints. The polymer derived from styrene, $C_6H_5CH:CH_2$, is used in household moulded goods, while vinylidene chloride, $CH_2:CCl_2$, forms a crystalline polymer used for spinning into fibres.

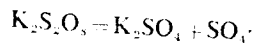
The physical forms of monomers vary: most of them are liquids, but some, including vinyl chloride (B.P. $-14^\circ C$) and butadiene, $CH_2:CH.CH:CH_2$ (B.P. $-3^\circ C$) are gases.

The Polymerization Reaction

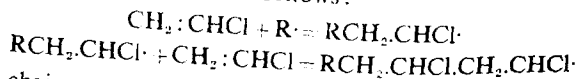
The polymerization reaction may be illustrated by reference to vinyl acetate:



In order to form the polymer a catalyst, or, more correctly, an initiator is necessary. This is a material that on decomposition generates *free radicals*, ephemeral chemical entities with only part of their functional groups satisfied. Free radicals are extremely reactive, and can react with the double bond of an ethylenic or vinyl compound, thus instigating a chain reaction. To be useful they must be formed at some temperature between, say, 40° C and 100° C. A simple example is the decomposition of potassium persulphate.



Another well-known decomposition is that of benzoyl peroxide to give $\text{C}_6\text{H}_5\cdot$. Polymerization takes place when the radical attacks the monomer. If the free radical is denoted by $\text{R}\cdot$ a simple scheme would be as follows:



A chain reaction is thus set up and continues until a secondary reaction occurs which terminates the chain. This may be caused by the combination of two free radicals, which mutually satisfy each other, or by a number of other side reactions.

An interesting development of the free radical catalysis of polymerization is reduction activation. This was discovered independently in Great Britain, the United States, and Germany about 1939-40. If a small quantity of a reducing agent such as sodium sulphite Na_2SO_3 , or sodium thiosulphate, $\text{Na}_2\text{S}_2\text{O}_3$, is mixed with a persulphate, it is found that the rate of polymerization is about 50 times as fast as with persulphate alone. This is because the reducing agents help in the generation of free radicals. A large number of other systems of this type are known.

The number of molecules that can combine to form a polymer may vary from about 20 to more than 10,000, the molecules in a polymer not all having the same molecular weight. Polymers made from a homologous series of vinyl monomers show a continuous gradation of properties. Thus polymethyl methacrylate is a very hard resin, polyethyl methacrylate is

hard, but can be softened more readily, while polybutyl methacrylates become quite tacky on slight warming. Higher poly-methacrylate, such as polylauryl methacrylate, are quite tacky at ambient temperature, but if the alcohol chain is made still longer as with polystearyl methacrylate, the resin takes on wax-like properties.

A brief reference might be made here to copolymers. It is possible for several different monomers to polymerize together to form a copolymer as distinct from a homopolymer, which is derived from one monomer only. Rates of copolymerization vary considerably, and sometimes only a small quantity of one monomer can be readily copolymerized with another. Thus commercial copolymers of vinyl chloride and vinyl acetate contain only about 13 per cent of vinyl acetate. This copolymer is softer and dissolves rather more readily than a homopolymer of vinyl chloride. On the other hand styrene and vinyl acetate cannot be made to copolymerize under any known conditions.

Frequently polymers are 'tailor-made', i.e., varied quantities of monomers are copolymerized so as to give specified properties. Thus vinyl acetate and ethyl acrylate can be copolymerized to give a flexible and extensible resin with good low-temperature flexibility.*

Emulsion Polymerization

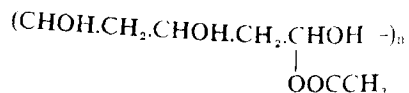
An emulsion of a monomer is a typical example of the system described on page 88. For simplicity, the droplets of monomer that are dispersed in the water phase may be considered to be spheres of diameter between 0.1 and 2.0 microns. When the monomer is gaseous, it is liquefied under pressure.

In order to obtain a stable polymer emulsion it is necessary to carry out the polymerization in such a manner that the emulsion remains stable throughout a complex chemical reaction, and that no precipitation occurs, because 'precipitation' implies the breakdown of the emulsion into water phase and solid resin. This latter either forms a solid lump which may contain some unreacted monomer, or else it may separate into discrete solid

* M. Leitner: 'Large Molecules', *Science News* 32.

particles, which, if spherical, may be up to several millimetres in diameter. It is therefore essential that the emulsifier system be properly chosen, although to a large extent the choice is empirical, and may be modified to suit the end use of the product. The polymerization conditions must be controlled so that the reaction is a steady one, because a violent exothermic reaction usually results in partial or even complete precipitation. Depending on conditions, cationic surfactants may be used instead of the more conventional anionic types. In general the percentage of emulsifier to monomer weight may be as low as 0.1 per cent when colloidal stabilizers are present; similarly, the percentage of colloid to monomer may vary from 1 to 6 per cent.

In addition to naturally occurring colloids, semi-synthetics, such as carboxymethyl cellulose, are often included in emulsions, but it is rare to find methyl cellulose used because it is not soluble in water above 35° C. Water soluble resins are frequently encountered. An example is polyvinyl alcohol, which is formed by partial hydrolysis of polyvinyl acetate and may be represented as:



Other well-known stabilizers are the salts of polymers of acrylic acid, $\text{CH}_2=\text{CH}\cdot\text{COOH}$, and methacrylic acid $\text{CH}_2=\text{C}(\text{CH}_3)\cdot\text{COOH}$.

The free radical catalyst used to initiate the polymerization reaction is normally water soluble, e.g., potassium persulphate or hydrogen peroxide. There is often a period of induction, varying from a few minutes to an hour or more, before polymerization starts, owing to the presence of traces of inhibitors which are slowly removed by the catalyst. Hydroquinone $\text{C}_6\text{H}_4(\text{OH})_2$ is a typical inhibitor; it is often deliberately added to vinyl monomers to prevent spontaneous polymerization. Sometimes a 'modifier' such as lauryl mercaptan, $\text{C}_{12}\text{H}_{25}\text{SH}$, which has the effect of reducing the chain length of the polymer formed, is included.

Industrial Polymerization Processes

A typical large-scale reactor is a jacketed stainless steel kettle of between 250 and 1,000 gallons capacity. It is equipped with a stirrer, and can either be operated with a reflux condenser or under autogenous pressure. Monomer and water are pumped in, emulsifiers being either predissolved, or added to the vessel. In some processes monomer is added gradually during the reaction, which takes from 1.5 to 6 hours, but in others the entire quantity is added at the start. Most polymerization reactions start between 40 and 80° C. depending on the monomer, the catalyst employed, and the quantity used. The reaction itself generates considerable heat, so control is effected either by efficient cooling of the outside of the vessel, or by the reflux. Stirring must be very carefully controlled, because too great a speed may precipitate a partially precipitated monomer owing to centrifugal action, while inefficient stirring may result in local overheating, and polymerization with almost explosive violence. The finished emulsion, in which there is only a very small quantity of residual monomer, usually under 0.5 per cent of the starting material, can be run out of the base of the reaction kettle, and is filtered from fortuitous aggregates that may have formed during polymerization.

THEORETICAL CONSIDERATION OF EMULSION POLYMERIZATION

Until recently the technology of emulsion polymerization far outstripped theory. In part this was because at first emulsion polymerization was mainly used in the manufacturing of synthetic rubbers, where there was a tendency to treat emulsion formulae purely as empirical recipes. In recent years, however, considerable theoretical progress has been made, as will now be explained.

Soap solutions have a number of anomalous properties compared with other salts of electrolytes. These include an abnormal solubility increase over a small temperature range, and abnormally low osmotic pressure accompanied by relatively high

electrical conductivity. McBain, who made a systematic study of the colloidal properties of soap solutions, suggested that these anomalies could be explained in terms of the formation of aggregates of soap ions termed micelles. Although there has been much discussion on the nature of micelles, these aggregates of soap molecules may be assumed to be oriented so that they are aligned parallel to each other. They are capable of pulling insoluble liquids into solution, the liquids being fitted between the hydrocarbon heads of the emulsifier in a sort of molecular sandwich. Thus, by the addition of two to three per cent of emulsifier, the solubility of styrene in water can be increased from about 0.003 per cent to about 0.1 per cent at 40° C.

In most cases of emulsion polymerization where a water-soluble initiator is used, it is considered that polymerization is initiated inside the micelles. Once initiated the polymer that is formed first absorbs more monomer by diffusion from the liquid, and also some of the emulsifier, to form a stable particle. Where polymerization is initiated by free radicals in the water layer, the micelles form a reservoir for the propagation. The picture is slightly different where an oil-soluble catalyst is used, when initiation must take place in the particles themselves, or where the monomer is appreciably soluble in water, as is vinyl acetate or acrylonitrile. The picture is, however, similar in that the emulsifier and colloid must stabilize both polymer and monomer particles.

Properties of Synthetic Resin Emulsions

Polymer emulsions are milky liquids containing between 25 and 65 per cent polymer. The viscosity of even the most concentrated emulsions is usually quite low because viscosity does not depend on the concentration of the resin. On the other hand some latexes prepared with polyvinyl alcohol are almost paste-like in consistency.

Resin emulsions have a very wide range of applications, and their most useful properties are based on their being formed in an aqueous medium. Apart from water being the cheapest solvent, it is non-toxic and non-flammable. In contrast to emul-

sions, the viscosity of a resin solution depends on molecular weight, although the relation is not a simple one. Solution viscosity also depends on the nature of the solvent, and it is only possible to have a concentrated solution of a low molecular weight polymer in a solvent such as acetone. Thus the highest normal concentration of a nitrocellulose lacquer is about 15 per cent. For many applications, however, it is desirable to have as concentrated a polymer as possible, and an emulsion, which may contain as much as 65 per cent of solids, proves a highly suitable method of handling resins. Since the loss of large quantities of solvent is avoided, as when spraying, it is unnecessary to install an extensive solvent recovery plant.

Polymer emulsions may be treated in the same way as lacquers or solutions, e.g., they can be pigmented to give surface coatings suitable for emulsion paints. There is, of course, the slight disadvantage that the presence, of necessity, of emulsifiers and colloids may detract from the qualities of the polymer, and they may also impair the water resistance of the resin derived from the emulsion. These difficulties can, however, be overcome to a large extent by careful formulation.

Many applications of resin emulsions depend on the nature of the films deposited when the emulsion dries. These films and their properties will therefore be considered next.

Emulsion Films

If a thin film of emulsion is allowed to evaporate it may form a continuous film of polymer. Thus, on evaporation, polyethyl acrylate emulsions deposit a continuous, elastic, and rather sticky film. In other cases, such as emulsions of polystyrene or polyvinyl chloride, evaporation leads to the formation of a powder. An intermediate case is that of polyvinyl acetate, or polybutyl methacrylate, which may deposit weak, crazed films, or firmer, slightly brittle, more continuous films depending on conditions of temperature, molecular weight of the resin, and to some extent the nature of the emulsifying agent.

There is a temperature known as the second order transition point (T_g or T_m) below which the properties of a resin change,

and it acquires to some degree the properties of a solid rather than a liquid. Above this temperature, rotation of polymer segments begins, and the resin acquires rubbery properties rather than those of a glassy solid. For an emulsion to produce a continuous film, T_g must be reduced below the drying temperature. Thus an emulsion of polyethyl acrylate ($\text{CH}_2\text{CHCOOC}_2\text{H}_5$)_n for which T_g is -20°C , will always dry at room temperature to give a flexible sticky film.

To alter the properties of polymers and make them softer and more flexible, plasticizers are added. Plasticizers are frequently, but not invariably, organic esters, which form solid solutions with polymers, and the plasticized resin may contain from as little as 2 per cent to about 60 per cent of plasticizer. These plasticizers often have a specific action, and not every ester is compatible with every resin. Dibutyl phthalate is probably the most widely used plasticizer, but dioctyl phthalate, tritolyl phosphate, and esters of adipic and sebacic acids are frequently encountered. To quote a typical example, methyl methacrylate, which deposits a white powder from emulsions, produces a clear flexible film when plasticized to a degree of 40 per cent (on overall film weight).

For polyvinyl acetate T_g is about 28°C , and it needs about 7 per cent of plasticizer to give an emulsion that will dry out as a continuous film at ambient temperatures. For polystyrene T_g is 83°C , and about 28 per cent of dibutyl phthalate is required to give a continuous film. Polymethyl methacrylate, containing 40 per cent of dibutyl phthalate, will deposit a continuous film that has found some use as a leather finish.

With some emulsions the film formed has a matt, translucent appearance; others are clear and glassy. To some degree it is a function of the nature of the resin, but other factors play a considerable part. These include the ease of 'flow-down' or integration of the film, but the main factor is compatibility of surfactants and colloids with the resin. In formulating an emulsion so that a clear glassy polymer film may be obtained, it is desirable that all emulsifying and stabilizing agents should form a homogeneous solid solution with the resin. Many surfactants

of the sulphate and sulphonate type are universally compatible with polymers, but on the other hand fatty acids, formed by loss of ammonia from ammonium soaps, are often incompatible, and they make resin films cloudy, or else form a greasy layer on the surface. A polyvinyl alcohol that contains no residual acetate groups gives a cloudy film with polyvinyl acetate.

PIGMENTED EMULSIONS AND THEIR FILMS

It has already been pointed out that one of the largest applications of synthetic resin emulsions is as a base for emulsion paints. We shall therefore consider next the preparation and properties of pigmented emulsions and their films. In this section the terms 'pigmented emulsion' and 'paint' will be treated as synonymous.

Formulation of a paint has to be made with great care to avoid precipitation and to give such desirable properties as long-term stability, and good brushability. For an emulsion paint it is usual to grind the pigments in water with a little dispersing agent and then to mix the pigment paste with the resin emulsion. Colloids such as methyl cellulose and carboxymethyl cellulose are often added to improve the general properties of the system. Sodium polymetaphosphate, which is sometimes used as a water softener, is often included, as well as preservatives to prevent microbial attack on the natural products that are present.

The mechanism of film formation when an emulsion film dries is still imperfectly understood, as is the more complicated system of film formation from an emulsion paint. Recent work, as illustrated by electron microscope and phase contrast photomicrographs, indicates that a touch-dry emulsion film is a close-packed array of discrete particles. These tend to flow together, the degree of integration depending both on the temperature and on the nature of the emulsifiers that are associated with the emulsion particles. If the film is below the second order transition temperature integration will not take place. (See Insets 11 and 12.)

The nature of a pigmented film depends to a considerable extent on the ratio of pigment to resin. This may vary from

about 3 parts by weight of pigment to 1 of resin (including plasticizer) where a flat finish is required, to equal parts of pigment and resin when an eggshell finish is preferred.

Electron micrographs indicate that in a freshly dried paint the polymer particles still retain their identity. Pigment particles on the surface are only loosely held, but within the film pigment particles tend to cluster. The electron micrographs reveal that changes take place in the film continuously over a twenty-week period, although the nature of the changes on ageing are still obscure. In some cases a clear layer develops on both the top and the under surface of the film.

Emulsion Paints

In Great Britain most emulsion paints are based on polyvinyl acetate, and increasing quantities are being used in the United States. Vinyl acetate monomer is cheap; it can be polymerized easily to give an emulsion containing 55 per cent or more solids, and when plasticized the polymer emulsions have good film-forming properties. They are resistant to alkali, unaffected by light, and tend to 'breathe', an advantage in many applications. The plasticizer most commonly used is dibutyl phthalate, although dibutyl glycol phthalate, which is less volatile, is sometimes also used. On the Continent, dibutyl phthalate is often blended with tritolyl phosphate, which has the advantage of imparting some resistance to fire. High molecular weight plasticizers, such as polyglycol sebacates, with molecular weights of the order of 1,000 and upwards, are sometimes included because they have less tendency to migrate into a substrate. The plasticizer content should be between about 12 and 20 per cent, depending on pigment loading.

Nowadays there is a tendency to utilize a copolymer emulsion, instead of one 'externally' plasticized with, for example, dibutyl phthalate. A co-monomer is selected which gives additional flexibility to the resin films; this is frequently referred to as an 'internal' plasticizer, although, in fact, it forms an integral part of the polymer molecule. Co-monomers include di-octyl maleate, ethyl acrylate, and more recently higher vinyl esters

such as vinyl caprate and vinyl stearate. The copolymer emulsions are claimed to have a number of advantages, one of which is non-migration of the plasticizer, and easier 'wetting' of pigments by the emulsion.

Styrene, and styrene copolymers, have been used in the United States to a considerable extent in emulsion paints. The copolymer is frequently a styrene-butadiene latex of high styrene content not unlike synthetic rubber (GR—S) in composition.

Polystyrene and its copolymers are, however, less light-stable than polyvinyl esters and they tend to yellow on ageing. While the water resistance of polystyrene and its copolymers is very good, the preparation of concentrated emulsions is rather difficult and plasticization is somewhat of a problem. These two difficulties seem to result from an integral property of the styrene molecule, probably due to the shape of the styrene molecule itself or of the styrene units in a polymer. Dibutyl phthalate is amongst the few satisfactory plasticizers for polystyrene, and it should comprise between 28 and 35 per cent of film weight. If the quantity of the plasticizer is raised, the resin becomes sticky; if less is present the resin film is 'cheesy' and without either flexibility or strength.

Emulsions of polyacrylates and some higher polymethacrylates, such as butyl polymethacrylate, are finding increased applications as bases for emulsion paints, but, because of their high price, mainly for specialized applications. They require no plasticizers; in general their sensitivity to water is much less than that of polyvinyl acetate, and the films are highly elastic, which tends to prevent cracking.

An entirely different type of emulsion used as a medium for emulsion paints is that based on alkyd resins. Alkyds are formed by condensing a polyhydric alcohol such as glycerol $\text{CH}_2\text{OH}.\text{CHOH}.\text{CH}_2\text{OH}$, with phthalic anhydride, $\text{C}_6\text{H}_4(\text{CO})_2\text{O}$, and a fatty acid, usually of the drying oil type derived from linseed oil. Alkyds, which are normally available in solvents such as xylene, dry by oxidation in the same way as linseed oil. Between alkyd and vinyl resin emulsions there is a fundamental

difference. It is not possible to prepare an alkyd resin by emulsion polymerization. The resin must be pre-formed and then converted into an emulsion. Ammonium soaps, the ammonium salt of casein, and sometimes non-ionic surface active agents are used as emulsifying agents.

Paints derived from alkyd emulsions are claimed to have better washability than those prepared from polyvinyl acetate, and adhesion to many surfaces such as old hard-gloss oil paints is claimed to be superior. This is not the case, however, with highly alkaline surfaces like fresh plaster because the alkali tends to hydrolyse the alkyd. The resistance of paint films based on alkyds to ultraviolet light is also inferior to that of paint films based on vinyl resins.

Emulsion paints are generally based on a tinted white pigment, the most reliable white pigment being titanium dioxide, which is very inert in either of its commercial forms anatase or rutile. Rutile is resistant to chalking, but has a slightly more yellow tinge than anatase.

Stability to Freezing

Some synthetic resin latexes can be frozen solid and afterwards thawed out without the emulsion being affected. With others an irreversible precipitation occurs on freezing and subsequently thawing. Where, therefore, emulsions or emulsion paints have to be stored during the winter months it is obviously necessary for formulations to be used that will withstand freezing. Research has shown that a number of factors are involved. Where freezing is accompanied by an uneven distribution of emulsifiers and stabilizers, the emulsion will probably be unstable to freezing. On the other hand reduction of particle size also decreases stability to freezing, probably owing to increased adsorption of emulsifier on the greater surface of the emulsion particles. Freeze-thaw stability also decreases as plasticizer content increases. This seems due to the increased tendency of the plasticized resin to flow together at cold temperatures by reducing the second order transition point. Where films are pigmented the tendency of the pigmented system or paint to be ruined by

freezing is slightly reduced, but it is desirable to use efficient wetting agents with the pigment. The addition of a water soluble antifreeze such as ethylene glycol has only a very slight effect in improving freeze-thaw resistance of pigmented emulsions.

OTHER SURFACE COATING APPLICATIONS OF RESIN EMULSIONS

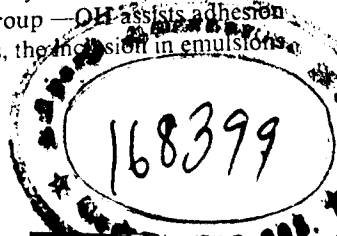
Leather finishing is one of the oldest and most conservative of crafts. Casein finishes have long been used, but in recent years a number of resin emulsions have found increasing application. The various polyacrylates, with their almost rubber-like properties, have proved ideal, but some polymethacrylates such as the butyl and ethoxyethyl esters have also found considerable use. Absence of tack on plating and a low temperature of cold crack ($< -5^{\circ}\text{C}$) are a desirable feature of this type of emulsion. Ethoxyethyl methacrylate emulsions are also used as a final dressing for leather because they give bright glossy films.

For special purposes paper can be coated with synthetic resin emulsions, although solutions in organic solvents are often used because of the tendency of some papers to ruck when water wet. Polyvinyl acetate is used to coat wrappings for foodstuffs, and paraffin wax may be added to increase imperviousness to water. This resin has also an interesting application in the direct coating of foodstuffs, the emulsion being applied by spraying or dipping.

Adhesives

A major application of resin emulsions is in the adhesive field. The theory of adhesion has been considered by de Bruyne in *Science News* 41. Although de Bruyne describes several simple principles which assist in assessing the potentialities of a resin as an adhesive, applications of resin adhesives are as yet to a considerable extent *ad hoc*. It must not be forgotten that adhesion is a function of both the adhesive and the adherend. Thus a good wood adhesive may be totally unsuitable for metals. Since the presence of the hydroxy group —OH assists adhesion considerably on many polar surfaces, the inclusion in emulsions

Doc 50146
100.46



of a partially hydrolysed polyvinyl acetate is beneficial; in fact, polyvinyl acetate emulsions, with partially hydrolysed polyvinyl acetate as emulsifying agent, are probably among the most widely used vinyl resin adhesives. They can be applied to the bonding of wood, paper, and cellulose products, most textile materials, and even for bonding certain materials to metals. In many cases, e.g., soft-wood glues consisting of polyvinyl acetate, the joint strength of the bond exceeds the shear strength of the wood itself, approaching 3,000 p.s.i. For reasons that have already been given, a polyvinyl acetate emulsion adhesive should be plasticized to the extent that T_g is below the temperature of drying. This is especially so when the adhesive is for flexible articles.

It will be noted that reference has been made almost exclusively to polyvinyl acetate and its copolymers, but the industrial use of synthetic resin emulsions as adhesives is by no means confined to this resin. Apart from the fact that its properties make it the most versatile, polyvinyl acetate has also the advantage of being the cheapest vinyl resin available. Among its other desirable properties, its emulsions are compatible with natural rubber latex, provided that the pH of the emulsion had been brought up to about 9.0.

Heat sealing is a method of forming a bond for which many resins, including acrylates and methacrylates, as well as the polyvinyl esters, are suitable. The surfaces to be sealed are coated with emulsion which is air dried; heat and pressure are then applied to form the bond. Cellophane can be stuck to itself or to paper by this method.

Textile Applications

The improvement of natural textile materials, as well as the development of new synthetic fibres, is one of the most important of industries. The treatments to which fibres and fabrics are subjected can be divided into two main groups. These include sizing, which is a temporary treatment intended to improve the strength of the fibre while it is on the loom, and finishing, i.e., processes carried out to improve the appearance of the

final piece goods or fabrics. Finishing can be reasonably 'permanent', but only a 'semi-permanent' finish is often all that is required. Typical examples of the latter are dilute polyvinyl acetate emulsions that are marketed to give a finish resembling starch, which will, however, withstand about five or six washings.

When the quantity of resin applied is very high, as in a leathercloth or a waterproofed fabric, the nature of the fabric is of secondary interest. Leathercloth, for which the traditional coating was nitrocellulose, is often now made from polyvinyl chloride or its copolymers. Although pastes, i.e., dispersions of resin in plasticizer, are generally used, it is possible to use emulsions. In order to obtain a tough flexible continuous film this must be dried at a high temperature, since polyvinyl chloride does not form a true plasticized film in the cold. Plasticized polyvinyl chloride film is probably a two-phase system consisting of a solid solution of the plasticizer in the resin, and a liquid phase consisting of a solution of resin in plasticizer. This equilibrium cannot be obtained, however, until the plasticized mix is heated.

Sizes for textiles consist mainly of water-soluble materials. As such, one would not expect emulsions to be used, but there are several references in the literature to the use of a polyvinyl acetate emulsion as a size for rayon, the resin being scoured off the fabric after weaving.

Many 'finishing' processes on textiles involve some degree of stiffening. A wide range of methacrylate and vinyl esters and their copolymers are claimed to be useful for this purpose. Sometimes they are added as urea-formaldehyde resins as in the well-known crease-resistant process. Polyvinyl acetate is used together with rubber latex in the manufacture of non-woven fabrics. A claim has been made that a copolymer of butadiene and a methacrylate reduces the shrinkage of wool. An emulsion composed of bromoform (CHBr_3) and triallyl phosphate ($\text{CH}_2=\text{CH}.\text{CH}_2$), PO_4 , is used for fireproofing.

Monomers can be converted into polymers after having been emulsified with a cationic emulsifier in exactly the same manner

as with the more conventional soaps and colloids. Few colloids are stable in the presence of cationic emulsifiers, but under some conditions polyvinyl alcohol can be used, especially where an emulsion is prepared by inversion. In this process a normal anionic emulsion is prepared and then reversed by controlled addition of a cationic reagent. Great care must be taken at the 'neutral' point to avoid premature precipitation.

A textile material has generally a weak negative charge, and when a dilute cationic emulsion is brought into contact with the textile the resin exhausts rather like a dyestuff. For this purpose the concentration of the resin is about 0.1 per cent: the aim is to obtain complete exhaustion with a resin uptake of about 2 per cent of the textile weight.

This treatment is frequently employed with nylon hosiery where the aim is to impart firmness, body, and snag resistance to make handling and packing easier. Polymethyl methacrylate, by virtue of the fact that it is a hard brittle resin, is suitable for this purpose. Positively charged polyvinyl acetate, when employed in a similar fashion, gives a softer and fuller 'handle' to the goods. Cationic emulsions of polystyrene and polymethyl methacrylate are also used for delustring textiles.

Cement Mixes

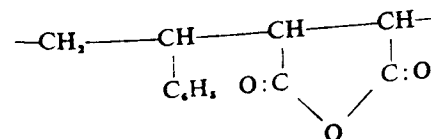
One of the most surprising uses of polyvinyl acetate emulsion is as an additive to cement. The function of the resin is to provide greater resilience and impart a warmer feel to the cement. The wetting agents present often affect the setting time of the cement, non-ionic surfactants tending to lengthen the time of set. The addition of resin increases the cost of the concrete very considerably so the quantity added must be kept as low as possible. A ratio of 1 part of cement to 0.2 of resin is the optimum, but less than a third of this concentration brings about useful modifications to the basic properties of the concrete.

ALKALI SOLUBLE EMULSIONS

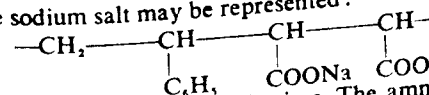
When a copolymer that contains an acid group is made, it is often found that the resin is alkali soluble, because the acid

Resins in Emulsion

group readily forms a salt soluble in water. Thus a single unit of a copolymer of styrene and maleic anhydride may be represented:



This polymer reacts with dilute sodium or potassium hydroxide, ammonia, and many organic bases, to give a clear solution. A unit of the sodium salt may be represented:



These solutions are used as textile sizes. The ammonium salts of these copolymers tend to precipitate on standing, partly owing to loss of ammonia with regeneration of the original acid, and partly owing to amide formation. Thus the ammonium salt of the copolymer of styrene and maleic anhydride can be used as a 'loom finish' for textiles, i.e., it is first used as a size which is then rendered insoluble. A logical development that has resulted for the preparation of alkali-soluble resins is that of alkali-soluble emulsions. Unsaturated acids, such as maleic and acrylic acids, are typical of those used for copolymerization. These acids are water soluble, while most monomers are virtually insoluble though slightly solubilized by the emulsifying agents. In consequence, owing to the lack of homogeneity of the system it is difficult to obtain an even copolymerization. However, the technological difficulties of this copolymerization have been overcome, and a number of alkali soluble copolymers are now available in emulsion form both in Britain and in the United States.

The copolymers of this type need only a comparatively small amount of alkali to dissolve them, about 10 per cent of the resin weight. The water-soluble resin solutions are extremely good binders for pigments, and have potential uses in paper coating and as thickeners. They are immune to damage on freezing because on thawing they are still water soluble.

CONCLUSION

The development of synthetic resin emulsions is an interesting example of the blending of two fields of chemistry to produce a valuable group of new products. It is also an example of the development of chemical theory in parallel with technology to the mutual advantage of both.

There is still much scope for progress, because the remarkable property of synthetic resin emulsions of providing a water-resistant coating that can be formed by evaporating an aqueous medium is capable of further industrial exploitation. These products will probably find an ever-increasing outlet in industry, paralleled by developments in the complex chemistry of resin emulsion systems. The author expresses his thanks to J. E. O. Mayne, D.Sc., F.R.I.C., for his valuable advice and assistance.

FURTHER READING

General

Leitner, M.: 'Large Molecules'. *Science News* 1954 32 43.

Couzens, F. G. and Yarsley, V. E.: *Plastics in the Service of Man*.

Emulsions

Clayton, W. (revised by Sumner, C. G.): *Emulsions and their Technical Treatment* (London, Churchill, 1954, 5th edition).

Ward, A. G.: *Colloids, their Properties and Applications* (London, Blackie, 1945).

Vinyl Resins

Morrel, R. S. and Langton, H. M. (Editors): *Synthetic Resins and Allied Plastics* (Oxford University Press, 1951).

Schildknecht, C. E.: *Vinyl and Related Polymers* (John Wiley, 1952).

RESEARCH REPORT

A. W. HASLETT

SICKLE CELL ANAEMIA

SICKLE cell anaemia is an inherited disease of the blood which, as a result of research from about 1948 onwards, has now an assured place in the histories of biochemistry, genetics, and anthropology. The most recent and spectacular contribution is from a biochemist – Dr V. M. Ingram of the Medical Research Council's unit at the Cavendish Laboratory, Cambridge. By a new combination of methods he has shown that the haemoglobin of a sickle cell patient differs from the common haemoglobin of human adults only by the substitution of one amino-acid for another in a chain of nearly 300 amino-acids. This is thought to be the closest approach yet made to identifying the chemical action of a gene.

Sickle cell anaemia has been recognized as a disease since 1917. The red cells of those who suffer from it have an abnormal haemoglobin – haemoglobin S – which, in the absence of oxygen, becomes insoluble in the cell fluid. It therefore crystallizes and distorts the shape of the cell. A variety of cell shapes are produced, one of which, that of a sickle or crescent, accounts for the name given to the disease. Few individuals, with only this abnormal haemoglobin in their red cells, survive to adult life. Those who inherit the abnormality from only one parent have a mixture of haemoglobins, and are said to show the sickle cell trait – as distinct from the disease. In them, the abnormality is harmless in any circumstances, and in one condition is an advantage. This is in a district where malignant tertian malaria is prevalent, red cells containing the abnormal haemoglobin being less well suited to the malaria parasite. Although there are points of uncertainty still to be cleared up, this is probably the clearest case

in which a gene that is lethal when inherited from both parents remains reasonably common none the less, because of an advantage that it confers when inherited from one parent. This is the interest of the disease to the geneticist. For the anthropologist, there is the interest that, with malaria control now extending, it should be possible to watch the gradual disappearance of an abnormality which in other circumstances might have been taken to be racial. At present, in tropical African tribes, up to 20 per cent of individuals, or even 40 per cent, show the sickle cell trait; in American negroes the proportion is about 8 per cent; in Southern Europe and Asia the abnormality is rare.

Dr Ingram's research goes back to two independent investigations in the United States – one genetic and the other chemical – which were published in 1949. (An account of this work was given in *Science News* 17, pp. 127–8; but the substance of it will be repeated here for convenience.) Dr J. V. Neel, of the Laboratory of Vertebrate Biology, in the University of Michigan, established the genetic basis of the disease. He tested for sickling the blood of one or, where possible, of both parents of patients with sickle cell disease. Where both parents were tested – which was in 13 cases – both of them showed the trait. In a further 16 cases, only one parent could be tested but showed the trait in all cases. Sickle cell anaemia was therefore an inherited disease determined by a single pair of alternative genes.

The chemical investigation, published soon after, was by Professor Linus Pauling and others at the California Institute of Technology. They prepared solutions of haemoglobin from normal people, those with sickle cell anaemia, and those with sickle cell trait. They then compared these solutions by the method of electrical separation known as electrophoresis. They found that haemoglobins of normal people and those with the disease were different, and that the haemoglobins of people showing the trait was a mixture of the two kinds. Sickle cell anaemia, they concluded, was an inherited 'molecular disease'.

This was the situation when Dr Ingram began his work, except that the difference between the haemoglobins had been shown to be chemically slight, and that further abnormal

haemoglobins had been discovered. Dr Ingram had therefore to use sensitive methods, and to apply them to specified haemoglobins – haemoglobin *A*, the common adult kind, and haemoglobin *S*, the abnormal sickle cell kind. He summarized the first stage of his research at the 1956 meeting of the British Association, when speaking on the structure of haemoglobins in general.

Dr Ingram described his new combination of methods as 'finger printing'. The combination was between electrophoresis – known already to show a difference – and the method of filter-paper chromatography. He began by breaking up the haemoglobin molecules into peptide fragments – chains formed by the linking together of amino-acids – with an average of about 10 amino-acids to each fragment. He did this with the digestive enzyme, trypsin. This is a standard procedure in protein chemistry; its effect on haemoglobin is to break up about 70 per cent of the molecule, the remaining 30 per cent being left intact. He then placed a drop of each mixture on one or other of two large squares of moist, thick filter-paper, near the middle of an edge, and applied an electric potential parallel to the edges on which the drops had been placed. This first separation depends, clearly, on differences in the charges carried by the fragments. He got a second separation by arranging for a solvent to be sucked up from the same edge. This was the chromatographic part of the separation, and it was in a direction at right angles to the electrical part. From each haemoglobin, the combined separation yielded 30 spots, of which 29 appeared to be identical, while the 30th spot was in different positions on the two sheets.

To pin down the difference precisely took more work than may have seemed probable at this stage, and the methods used do not lend themselves to detailed, short summary. It will be said, therefore, only that the fragment that showed a difference was found to be built up, in each case, of nine amino-acids; that the identities and positions of these in the chain were established; and that the difference amounted to the replacement of one unit of glutamic acid – in haemoglobin *A* – by a unit of valine in haemoglobin *S*. Valine is an amino-acid with one

amino and one acidic (carboxyl) group – the least possible for participation in a chain built up through amino-acidic linkages. Glutamic acid has an acidic side-chain; and the presence of an extra unit of glutamic acid therefore adds an additional negative charge to the fragment containing it.

The rest of Dr Ingram's work was confirmatory. In the case of the other 29 spots, he did qualitative analyses of the amino-acids that they contained – and found no apparent differences. There remained the major fragments – 30 per cent of the original molecules – which had been left untouched by trypsin. To break up these, he used another protein-digesting enzyme, chymotrypsin, and so obtained a further collection of fragments, which he again separated and studied for differences.

Dr Ingram quotes results of X-ray studies in Cambridge, and of unpublished chemical work by Professor Pauling, to support the conclusion that the complete haemoglobin molecule consists of two identical half-molecules. Adding this conclusion to his own work, the only difference in composition between the common adult haemoglobin and haemoglobin *S* is seen to be the replacement, in each half molecule, of an amino-acid with an extra negative charge by one that is neutral. Possibly there may be a difference also in the folding of the chain. Evidence to prove or disprove this does not exist.

According to present ideas, a chromosome is a long coiled filament of deoxyribose nucleic acid (DNA); a gene is a section of this filament; the first action of a gene is to mould from itself a corresponding section of ribonucleic acid (RNA), which passes out of the cell nucleus into the surrounding cytoplasm; and finally, the RNA serves in turn as a mould from which a protein – determined initially by the gene – is made. If this picture is correct, the identification of a single small difference between proteins determined by alternative genes, is probably as close as we may hope to get for some long time to watching a gene in action. The nearest previous approach came from some research on the biochemistry of anthocyanin flower pigments, done by Dr Rose Scott-Moncrieff at the John Innes Horticultural Institution in the middle 1930s in co-operation with Pro-

fessor (now Sir Robert) Robinson on the chemical side. But in this case there would be protein enzymes between gene and pigment changes, whereas haemoglobin is probably the first protein made by the gene.

Although the story of haemoglobin *S* seems now to be surprisingly complete, there is a complication in the background; at least seven further types of abnormal adult haemoglobins are known to exist. Two of these were discovered by Dr Neel or a colleague within two years of his original work. Their characteristic in terms of the disease is that a child who inherits haemoglobin *S* from one parent, and one of these further abnormal haemoglobins from the other, suffers from sickle cell anaemia in a less severe form. Both of them, like haemoglobin *S*, appear to be regional in distribution. One of them, haemoglobin *C*, has a frequency of up to 10 per cent in parts of West Africa, where it maintains itself alongside haemoglobins *A* and *S*; the other partly replaces haemoglobin *A* in northern India. A third has been found in south-east Asia, and there are a further four that are thought so far to be very rare. It has been suggested that all these haemoglobins are variations on a single theme; that a single gene, which can exist in a corresponding number of forms, is responsible for all of them. Dr Ingram has said that the chemistry of some of these further haemoglobins will now be studied. There is also room for research on the advantage that the three more common of them may confer on their owners, possibly in relation to types of malaria prevalent in the same regions.

A general introduction to the subject is given in J. B. S. Haldane's *Biochemistry of Genetics* (1954). The chief earlier papers referred to are: Neel, J. V., *Science*, 1949, 110, 64 and Pauling, L. and others, *Science*, 1949, 110, 543. The immediate references are: Ingram, V. M., *Nature*, 1956, 178, 792 and 1957, 180, 326.

NOBELIUM: ELEMENT 102

The making of element 102, for which the name 'nobelium' has been proposed, has been achieved lately by a joint Anglo-

Swedish-American research team, working with the 225-cm cyclotron of the Nobel Institute for Physics in Stockholm. The method used was the bombardment of curium-244, deposited on thin aluminium foil, with ions that included both carbon-13 and carbon-12 particles – those of carbon-13 being concerned in the wanted reaction. Curium being element 96, and carbon being element 6, the fusion of these two nuclei provided the number of protons needed to make element 102. On the basis that carbon-13 nuclei were the operative bombarding particles and were captured, the compound nucleus first formed would be nobelium-257. But, because of the energy of bombardment – which had to be not less than that needed to carry the carbon-13 nucleus through the potential barrier presented by the curium nucleus – the nucleus first formed would have an excess of energy, which would result in a number of neutrons being 'boiled off' from it. The question, then, was whether the number of neutrons lost in that way could be kept sufficiently low to leave a not too unstable isotope of the new element. But in this part of the periodic table, the probability of breakdown by spontaneous fission – as distinct from fission following absorption of a neutron – is high in suitable isotopes. It is highest in the so-called even-even isotopes – those in which the numbers of protons and neutrons are both even. For this reason, either nobelium-254 or nobelium-252 (152 and 150 neutrons respectively) would be unpromising isotopes to make because they would break down too quickly by fission. It seemed, however, that 3 or 4 neutrons was about the smallest number that it would be possible to boil off. It followed that nobelium-253, expected to break down by emitting an alpha-particle, otherwise a helium nucleus, was the first tolerably stable isotope that it was likely could be made in this way. The procedure for making it was therefore to keep the bombarding energy as low as possible, consistent with just penetrating the potential barrier. In that way it might be hoped to halt the boiling off process at 253, and avoid 252 which would break down by fission. Therefore, as the official statement expressed it, 'the experiment required very close control of the energy of the incident beam of bombarding

ions'. This was a new situation in the making of artificial elements, although one that had been foreseen.

The nobelium atoms formed were collected, along with other atoms produced and ejected by the interaction of bombarding particles and target, on a separate collecting foil placed close to the target. The collecting foil was effectively platinum. The atoms that were collected were washed off it with dilute acid, and the resulting solution passed down an ion exchange column. Finally, the nobelium was extracted from the column by an organic solvent. The identification of the new element depended on the association of a radioactivity close to that expected with this final solution. The half-life of the isotope made was about 10 minutes, and the energy of the alpha-particles emitted was about 8.5 MeV. The identity of the isotope has not yet been otherwise confirmed. By using more massive nuclei for bombardment, it appears probable that more massive and long-lived isotopes of nobelium could be made.

The members of the research team responsible were as follows: from the Nobel Institute for Physics, H. Atterling, physicist, W. Forsling and L. Holm, chemists, and B. Alström, electrical engineer; from the Atomic Energy Research Establishment, Harwell, John Milsted, a chemist who last year worked at the Argonne National Laboratory under exchange arrangements, and helped to separate the curium used as target material, and A. Beadle, a chemical technician; and from the Argonne National Laboratory, Paul R. Fields, a group leader in the chemistry division, and A. M. Friedman, now working at Harwell under exchange arrangements.

RADIO TELESCOPES

The bringing into operation of the 250-ft-diameter steerable radio telescope at the Jodrell Bank Experimental Station of Manchester University coincided within a few days with the opening of the Mullard Radio Astronomy Observatory of the Cavendish Laboratory, Cambridge. Questions of costs and administration are no part of the progress of science. However,

in view of criticisms which have been made publicly, it should be said about the Manchester telescope that the primary change in design was the substitution of a continuous metallic reflector for one built up of wires, so increasing both weight and wind resistance; that the reason for the change was the desirability of extending the wavelength range of the instrument so as to permit its use on centimetre wavelengths; that the scientific case for this extension became compelling only after work on the telescope had begun (see, for example, the article on 'Interstellar Material' by F. D. Kahn in *Science News* 34); and that the Department of Scientific and Industrial Research has stated recently that, in their view, an additional grant from Government funds would be justified. The Manchester radio telescope has performed well during tests. It is an all-purpose instrument; and the case for it is stronger in its final form than as first proposed.

The interest of the Mullard Radio Astronomy Observatory at Cambridge is that it carries to a new stage of development the principle of making use of the interference patterns that are formed when signals received in separated aerial systems are combined. To this a further principle has now been added – that of combining signals from a fixed and a movable aerial system, and making use of observations made on different days, with the movable system in successively different positions, to obtain the effect of an aerial array very much bigger than exists physically. The principle is described as 'aperture synthesis'. Clearly, the sources observed must emit steadily; otherwise the method would not be valid. In the observation of distant galaxies, and the general background of radiation from the Galaxy, this condition should be met.

The aerial systems used are similar to those used earlier at Cambridge. The rotation of the earth takes the place of one direction of aerial movement. Only one direction of movement is then necessary – that represented by dipping towards the northern or southern horizon – and the shape of the reflectors can be that of a half-cylinder. On metre wavelengths the reflectors can thus be built up of a comparatively small number of

wires, strained lengthwise along the cylinder. With the interferometer method of working, there is freedom to design separately for sensitivity and resolution; sensitivity being proportional to the combined area of the two aerial systems, and resolution to the distance between them. This means that the individual jobs to be done must be considered on their merits, and separate instruments built to suit them. Wavelength, once chosen, cannot be easily altered. With a cylindrical reflector the waves reflected are brought to a focus along a line, the axis of the cylinder, and many receiving elements are necessary. But in the case of an array built for a specific job, it is not a very serious restriction to have to choose a suitable wavelength and stick to it. There are two principal instruments at Cambridge. One of them operates on a wavelength of 1.7 metres, and is for work on extragalactic sources. The fixed array is 1,450 feet long and 60 feet wide; the movable array is 190 feet long and runs on railway tracks of that width over a distance of 1,000 feet. The collecting area is equivalent to about 190,000 square feet, nearly four times that of the Jodrell Bank steerable telescope, and in round figures a thousand times as great as that of the biggest optical telescope. With this instrument it should be possible to distinguish between directions separated by only a few minutes of arc, or less than a minute of arc in the case of a bright source. The other principal instrument will operate on a wavelength of 7.9 metres. Its use will be in a survey of the general background of radiation originating within the Galaxy. This, too, makes use of a movable array, but one of such light construction – weighing only about a hundredweight – that it can be picked up and moved from one position to the next by hand.

Ryle, M. *Nature*, 1957, **180**, 110.

HIGH-TEMPERATURE PLUTONIUM REACTOR

There are so many kinds of nuclear reactor already that a suggestion that any new kind is of special interest needs to be justified. Let it be said then that the new reactor will be fuelled with plutonium carbide; that fuel and moderator will be mixed in

ceramic-type tubes, through which gas under high pressure will be circulated; that it is hoped to dispense with canning material of any kind; and that operating temperatures of up to 800°C are to be aimed at. Such a reactor is the latest proposal for the development of the high-temperature, gas-cooled system that has been mentioned in the last two annual reports of the United Kingdom Atomic Energy Authority, and on various occasions by speakers connected with the Authority. The most recent developments are that portions of fuel assemblies of this kind have been made up and tested at Harwell, and that a contract has been given to the General Electric Company Ltd for much of the design and construction of a zero-energy system for experimental operation. This is to be built at the new Winfrith out-station of the Research Establishment.

Such a reactor would achieve two main objectives at once. First, it would provide an alternative to the Dounreay fast breeder reactor in the North of Scotland as a way of using as fuel the plutonium that will be produced in increasing quantities through the operation of reactors developed from the Calder Hall type. Secondly, the elimination both of metallic elements – which are subject to changes of physical state with temperature – and of canning will remove present limitations to working temperature, and should reduce costs. But such a reactor would not be able to breed – that is, to produce more fissile material than it consumed. The reason for this is that, whereas a fast reactor working on the plutonium/uranium-238 cycle produces – ideally – 2.6 neutrons for each neutron absorbed, in a thermal reactor (maintained by slowed-up neutrons), only 2.03 are produced. Since one neutron is required to maintain the chain reaction, and a second to replace each atom that has undergone fission, it follows that only the excess above 2.0 is available for breeding, and to cover inevitable losses from the absorption of neutrons in cooling and other materials.

For this reason, it was first proposed to develop the high-temperature, gas-cooled reactor for operation on the uranium-233/thorium cycle, since this is the more favourable of the two cycles for breeding in a thermal type of reactor. The fact that

plutonium will, within a few years, be becoming readily available as fuel has led to the decision to try the plutonium carbide version first. The purpose of the Winfrith assembly will be to test the nuclear properties of such a system at different temperatures. The effects of neutron irradiation on the fuel elements can be examined at different temperatures in the Harwell research reactors.

ICE NUCLEI AND RAIN

The theory that, in most conditions, raindrops in temperate climates owe their existence to the prior formation of ice crystals was first put forward by a Norwegian meteorologist, T. Bergeron, as long ago as 1933. It has been the basis of most modern attempts at rain-making; and these, in turn, have led to greater interest in the natural process. One of the problems is that crystals of ice – like crystals in general – find it difficult to make a start without some suitable small particle to grow on. In the case of water droplets, there is a gap of some 35 degrees between the temperature (-41°C) at which small droplets of pure water freeze spontaneously and that (sometimes as high as -6°C) at which ice crystals may begin to appear in clouds. There must thus be ice-forming nuclei in the atmosphere which are effective at temperatures only slightly below freezing-point. Research on the nature of these nuclei was presented by Dr B. J. Mason, of the Imperial College of Science, London, at the Dublin meeting of the British Association.

A suggestion that meteor dust might be important was put forward some years ago by Dr E. G. Bowen, of the Radio-physics Laboratory of the Australian Organization for Scientific and Industrial Research. His argument – which was contested – was based on apparent coincidences in time between the arrival of meteor streams in the earth's atmosphere and increased rainfall. Dr Mason's evidence is that, in laboratory experiments, particles of meteor dust do not appear to be effective ice-forming nuclei. Positively, he finds, first, that a number of silicate minerals such as kaolinite, microcline, and some micas are effective at temperatures higher than -15°C ; and, secondly,

that particles of the more common clay minerals, while initially active only at lower temperatures, when they have once functioned as ice-forming nuclei, are thereafter active at high temperatures. More precisely, if the ice crystals formed on such particles are allowed to evaporate in dry air at 0°C , the remaining clay particles are active in later experiments at temperatures as high as -5°C . Dr Mason's suggestion – so far unproved – is that their improved activity is due to the retention of small quantities of ice in submicroscopic cracks in the particles. The circulation of common clay particles to sufficient heights in the atmosphere might thus, he suggested, make the same particles effective at lower levels and temperatures in natural rain formation.

CORRESPONDENCE

THE ORIGIN OF THE 'BARN' UNIT

Dear Sirs,

In *Science News* No. 45 (page 106) T. M. Dunn repeats the physicists' tale that the name of the nuclear cross section unit 'barn' is due to the late Enrico Fermi. When I met Fermi in Chicago in March 1954, a few months before his most untimely death, I asked him about this story. He told me that the story is wrong. He knew very well that the invention of the name 'barn' was ascribed to him, but he assured me that he had never made the often-quoted utterance 'It's as big as a barn'. Fermi certainly was one of the greatest and one of the most amiable scientists of our age. I think it does not reduce his reputation if it is stated that the very amusing story about his introduction of the term 'barn' is untrue.

The riddle of how the rather unscientific word 'barn' was brought into grim nuclear physics is solved in the book *Sourcebook on Atomic Energy* by Samuel Glasstone (Van Nostrand, New York), where a footnote on page 264 says:

The term 'barn' was proposed in 1942 by the American physicists M. G. Holloway and C. P. Baker as the result of a broadly humorous association of ideas. It served the purpose of a code word, which was desirable at the time, and seemed appropriate because 'a cross section of 10^{-24} sq. cm. for nuclear processes was really as big as a barn'.

I should like to use this opportunity to tell you how much I appreciate *Science News*. I think I am one of the few Austrians (or the only one) who possesses all issues of *Science News* from the beginning. I owe much knowledge and many 'thrilling' hours to your excellent articles.

I also have a little wish. Please do not publish indexes so often because if you want to look something up you have to peruse

so many. A cumulative index of all 50 issues in *Science News* No. 50 would be most welcome. Although it would take much space it would be of great value to the reader. I, and I think many other readers, would be prepared to pay more for No. 50 if it contained a complete index of all previous issues.

Yours sincerely,

Friedrich Katscher

(DR FRIEDRICH KATSCHER)

Mariahilfer Strasse 133, Vienna xv,
Austria

17 Sept. 1957

The Editors would welcome readers' comments on the suggestion contained in Dr Katscher's letter - that a cumulative index to the first fifty volumes of Science News might be published either as part of No. 50 itself, or possibly as a special supplement.

SOME BOOKS RECEIVED

BRITAIN'S NATURE RESERVES by E. M. Nicholson (London, Country Life, 1957), pp. 175, 30s.

Nature Reserves in Britain are a relatively new development, and there are still many people who neither fully realize what they are, or what their function is. It is with this in mind that E. M. Nicholson, who is Director-General of the British Nature Conservancy, has written what amounts to a guide and handbook to Britain's fifty-odd National Nature Reserves managed by the Nature Conservancy, and other bodies, for the State. These range from the vast 40,000-acre reserve in the Cairngorm Mountains between the valleys of the Dee and the Spey, to small reserves of a few acres that have been set aside for special reasons. They include moorlands, woodlands, lakes, downlands, heaths, bogs, fens, sand dunes, coastal cliffs - all part of the living heritage to which, in general, we've paid so little regard in the ruthless past.

Before too late, however, the leadership and vision of a few pioneers like Charles Rothschild and Arthur Tansley (to whom the book is dedicated) led to the creation of these areas devoted to the study of nature - outdoor living museums in which the works of nature can be conserved as a national heritage, outdoor laboratories in which complex ecological studies can be carried out. And it must be emphasized that they are much more than preserves, they are true laboratories set aside for scientific research. Examples given by Nicholson are experiments on how native trees grow, and their influence on the soil, or on the effect of burning moorland, heavy rainfall, or grazing, on population dynamics.

And there is a point to be noted here. One does not barge into a scientist's laboratory uninvited, though one will almost always find a welcome if one takes the necessary preliminary steps and presents proper credentials. The same is true of nature reserves. While a number are freely open to visitors, others have delicate scientific experiments going on in them, or are there specifically for the protection of their wild life. One of the main objects of Mr Nicholson's book is to explain which reserves in Britain can be visited without formality, and what steps have to be taken to obtain permission to visit the others.

The volume is illustrated with over sixty photographs which reveal the extremely varied habitats that are available.

NOTHING LIKE SCIENCE by Magnus Pyke (London, John Murray, 1957), pp. 172, 16s.

It is not obvious whether the title of Dr Pyke's book is meant to disarm the critic of his wilder extravagances by disclaiming any close connexion with science, or whether it describes one of the articles of his faith. Be that as it may, not taken too seriously, *Nothing Like Science* is fun.

Should you have a scientist uncle who takes himself and his science just a little too seriously, then this, without doubt, is his Christmas present. But a word of warning. The volume is largely made up of reprints of Dr Pyke's highly amusing B.B.C. talks although there is no indication of this in the book itself, and Uncle may well recollect having heard them.

Witty and provocative though the talks may be, they do, like the story of Frankie and Johnnie, have a moral, which is simply that 'it is time we took a cool, level-headed look at the science of today, appraised its value, and the direction in which it is tending'.

A HISTORY OF INDUSTRIAL CHEMISTRY by F. Sherwood Taylor (London, Heinemann, 1957), pp. 467, 30s.

If there is one over-riding criticism of this, the last book from the pen of Dr Sherwood Taylor, it is that the author courageously set himself a more-than-human task which failing health prevented him from fully accomplishing, and it is sad to have to record that he only completed his manuscript shortly before he died.

The history of industrial chemistry, even when we apply the term in an extended sense to the empirical crafts of a pre-scientific age, is so much the history of material civilization itself (the author uses the term 'industrial chemistry' to include all modifications of the composition of matter that man has undertaken since the earliest times) that only the most brilliant mind in the full flux of its power can encompass the full sweep of the subject and its multitudinous ramifications. Technically the topics range from Chinese ceramics to artificial silk, from mediaeval glass to oil refining, from the fire-producing devices of the 16th century to the horrifying devastation

caused by the fission of uranium and plutonium, from medicine near allied to magic to the modern magic of the antibiotics.

To range so widely is virtually to attempt the impossible, and in the reviewer's view Dr Taylor would have left us a greater monument to his talents (and they were many) had he chosen to restrict himself and so achieve a more penetrating analysis of a narrower field. As it was, the author clearly recognized the magnitude of his self-set task, and in the absence of knowledge of the men and organizations that have created the chemical industry, he concentrated on chemical principles and the relation of industrial developments to chemistry itself. In the latter third of the book, therefore, there is a substantial substratum of purely theoretical chemistry and a corresponding lack of discussion of the ever-increasing impact of discoveries themselves on the industry, and of the industry on society.

Dr Taylor is at his best when he is writing about the long period that leads us to the rise of a scientific chemical industry; less so when he is trying to cram the multitude of developments of the past quarter of a century on to an already overcrowded canvas. Nevertheless, his book is a valuable addition to the scant histories of one of the greatest, both economically and socially, of man's manipulations of nature. There is a valuable list of suggestions for further reading and a good index. To Mr F. Greenaway, who saw Dr Taylor's manuscript through the press, we owe a special debt of gratitude.

SCIENTIFIC AMERICAN BOOKS (i) *The Physics and Chemistry of Life* (pp. 270); (ii) *Twentieth-Century Bestiary* (pp. 237); (iii) *Automatic Control* (pp. 148) (London, Bell, 1957), 13s 6d, 13s 6d, and 10s 6d respectively.

The *Scientific American* is known throughout the world for the excellence of its articles in which it interprets the advances of science to a public not necessarily professionally scientific. Now its editors have prepared groups of these articles, and have issued them in the above three volumes. Each volume covers one field of rapidly advancing knowledge.

'The Physics and Chemistry of Life' consists of some eighteen articles about life viewed as a physico-chemical process, and covering such topics as The Origin of Life, Photosynthesis, The Reproduction of Viruses, The Structure of Proteins, The Nerve Impulse, and the Electrical Activities of the Brain.

'Twentieth-Century Bestiary', despite its somewhat precious title

includes interpretations of such classic researches as Niko Tinbergen's on the courtship of animals, and D. and E. Lack's on their observations on the swifts in the tower of the University Museum at Oxford. In all there are twenty-four articles, which range from the Origin of Species to Physiological Engineering.

'Automatic Control' is a collection of eight articles on Feedback, the Second Industrial Revolution, Information, and the more philosophical 'Imitation of Life' by W. Grey Walter, and 'Man Viewed as a Machine' by J. C. Kemeny.

Anyone who values the *Scientific American* will be glad to have a selection of its invaluable articles available again in this useful integrated form.

THE WORLD OF THE SOIL by Sir E. John Russell (London, Collins, 1957), pp. 237, 25s.

This is the latest addition to the New Naturalist series: its author the doyen of soil scientists, who, at the age of eighty, has provided a volume full of the wisdom of his years. If any young scientist, surfeited by the content of his formal training, should feel that there are few fields left for him to explore, here is the book to dissipate his mental indigestion, because, time and again, Sir John, with the candour that comes from long experience, tells us, in reply to his self-set questions: 'We just do not know.'

Agricultural science, perhaps more than other branches of intellectual endeavour, has passed through phases at the end of which it truly appeared to those of the day that the whole truth was in sight. Such is the deceptiveness of partial knowledge; today it seems that the deeper the investigations of the soil scientists probe, the further recedes the horizon of ultimate truth.

For thirty years Sir John directed the premier agricultural research institution of the world – the Rothamsted Experimental Station at Harpenden in Hertfordshire. Now he has set out in broad outline the discoveries of those who have contributed to the revolutionary discoveries made about the soil in recent years – pedologists, chemists, physicists, biologists, and geologists – for in the soil is all science united. 'At Rothamsted' runs through the volume like a *leitmotif*, not that Sir John's outlook is in any sense parochial – indeed the planet is his parish – it is simply a measure of the institution and the renown his tenure of office raised it to.

Throughout the volume there are many passages that bear com-

parison with the best scientific writing that has come from the pen of any author. Here, for example, is how he introduces us to his subject:

'The soil is among Nature's greatest marvels. A clod of earth, seemingly simple and lifeless, is now known to be highly complex in structure, its particles most elaborate in their composition, with numerous invisible crevices inhabited by prodigious numbers of living organisms inconceivably small, leading lives of which we can form only the haziest conception yet somehow linked up with our lives in that they produce the food of plants which constitute our food, and remove from the soil substances that would be harmful to us. It would be flattering to our self-esteem to think that they are there for that purpose, but it is safer to take the wider view and to suppose that they, like ourselves, are leading their own lives, the purpose of which is completely hidden from us.'

The reviewer's only complaint about this fascinating volume is the unusual luxuriance of the crop of printers' errors, the product of slipshod proof-reading which extends to giving references as p. 000.

ABOUT OUR CONTRIBUTORS

D. M. ADAMS was born in 1933 and early developed an interest in things scientific. He was educated at Reigate Grammar School and King's College, London, where he graduated in 1954 with first class honours in chemistry. He received his Ph.D. for researches on the physical chemistry of bacteria and is now working on chemisorption at Imperial College, London. Apart from scientific studies his main interests are reading history and theology to complement his scientific work, and in walking when he cannot afford petrol.

P. E. HODGSON graduated in physics at Imperial College, London, in 1948, and subsequently studied the nuclear disintegrations caused by the cosmic radiation, for which he was awarded a Ph.D. in 1951. Afterwards he moved to University College, London, and began theoretical studies on nuclear collisions. He is now lecturer in physics in the University of Reading.

E. J. PRYOR, before entering the Royal School of Mines, served articles to Francis Sutton, the 'father' of volumetric analysis in Britain. This early training aroused his interest in applied science and later, with the discovery of the flotation process, was to forge for him the link between physical chemistry and the processing of ores. After twenty years in the Dominions he returned to his old college and became instrumental in organizing a course leading to the B.Sc. degree in Mineral Dressing. He re-designed the section of the Bessemer Laboratory now devoted to teaching and research in this expanding subject. What free times these activities leave him he fills with writing, gardening, and the enjoyment of the arts.

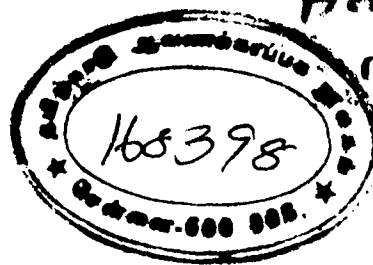
G. W. G. SHARP was born in 1929 and educated at Doncaster Grammar School and the University of Nottingham. He graduated B.Pharm. in 1949. He saw service in the Royal Navy from 1952 to 1954, during which period he studied Russian and is now on the translators panel of the Pergamon Institute which specializes in the translation of Russian scientific research papers. Was appointed lecturer in pharmacology in 1956 and is now engaged on research into biological rhythms concerned mainly with the adrenal cortex, and also upon the influence of light on the pituitary gland. In this connection in the summer of 1957 he led an expedition of members of the Nottingham University Exploration Society to Spitzbergen, where human 'guinea pigs' investigated the effects of light on

certain biological cycles in man. Outside the scientific field he is interested in sports, music, and mountains.

J. R. SINGER was born in Cleveland, Ohio, 1921. He was a navigating officer on troop transport ships during the war, spending about two years in the European theatre and two in the South Pacific area. Afterwards he went to the University of Illinois for a B.S. in mathematics, Northwestern University for his M.S. in physics, and the University of Connecticut for his Ph.D., where he worked with Professor Charles A. Reynolds, one of the important contributors to our understanding of superconductivity effects. He has taught at the University of Connecticut, the Catholic University of America, and George Washington University. He spent a year as research physicist at the Naval Ordnance Laboratory, and is at present chief staff physicist with the National Laboratories, Washington, D.C.

W. C. VAUGHAN M.B.E., is a Wing Commander in the Royal Air Force, stationed at present in Wiltshire. He was awarded the degree of Ph.D., for experimental work in radio physics in the University of London. He is the author of a number of scientific papers of which several have appeared in *Science News*.

H. WARSON was educated at the Central Foundation School, Cowper St., London, and Queen Mary College, University of London, where he graduated in 1938; he was awarded the F.R.I.C. in 1948. He is Research Manager at Vinyl Products, Ltd., Carshalton, Surrey, where he has been in charge of the research laboratory since 1945; he is responsible for a large number of patents concerning developments with synthetic resins, in particular vinyl resin emulsions. Since 1954 he has been honorary editor of *The British Chemist*, the journal of the British Association of Chemists.



SEEING CRYSTAL LATTICES AND DISLOCATIONS: *J. W. Mente*
THE STRUCTURE OF THE NEW TESTAMENT: *A. Q. Morton*
YEASTS AND THE BIOCHEMIST: *A. H. Rose*
THE FIRST AUSTRALIANS: *H. A. Lindsay*
LAND RESEARCH IN AUSTRALIA: *Archie Clow*
FOSSILS AND THEIR ENVIRONMENTS: *F. H. T. Rhodes*
MODELS OF THE ATOMIC NUCLEUS: *G. R. Satchler*
ALPHA, BETA, AND GAMMA RAY SPECTROMETERS: *I. R. Williams*

A NEW OUTLOOK IN ORGANIC CHEMISTRY: *R. S. Nyholm*
BLOOD: *M. C. H. Dodgson*
SPEECH ENGINEERING: *D. J. Bruce*
FIRST TRANSATLANTIC TELEPHONE CABLE: *Archie Clow*
NEW MAGNETIC MATERIALS—THE FERRITES: *F. Brailsford*
MALE OR FEMALE?: *D. B. Carlisle*
ULTRASONICS IN INDUSTRY: *A. Linford*

THE INTERNATIONAL GEOPHYSICAL YEAR 1957-8: *D. C. Martin*
 NUCLEAR COLLISIONS: *P. E. Hodgson*
 LOGIC AND BEHAVIOUR: *F. H. George*
 THE HISTORY OF THE MOON: *D. W. Allan*
 THE POLYSACCHARIDES OF HUMAN TISSUES: *S. A. Barker and*
M. Stacey
 THE NOT-SO-RARE EARTHS: *T. M. Dunn*
 RESEARCH REPORT: *A. W. Haslett*

[illegible]