

29

3198
94

Ocean-bed Prospecting
Movement of Ions in Metals
Methods of Technology
Chemical Technology
Generation
Phenomena of Cosmic Flight
Research Report

0

SCIENCE NEWS

EDITED BY A. W. HASLETT

29

PENGUIN BOOKS

MELBOURNE · LONDON · BALTIMORE

First published August 1953

*Am 1, 246
253-29*

Subscriptions to
SCIENCE NEWS

and

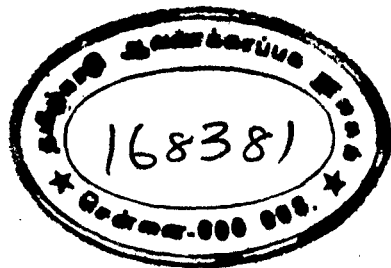
NEW BIOLOGY

are accepted at

9s. each

for four issues

including postage



*Made and printed in Great Britain
for Penguin Books Ltd
by The Campfield Press, St Albans
Collogravure Plates by
Harrison and Sons Ltd*

CONTENTS

OCEAN-BED PROSPECTING T. F. Gaskell and M. N. Hill	7
THE MOVEMENT OF INDIVIDUAL ATOMS IN METALS R. S. Barnes	32
METHODS IN PSYCHOLOGY F. H. George	45
CHEESE BIOLOGY James A. Barnett	60
GERMINATION A. M. Mayer	84
PHENOMENA OF SUPERSONIC FLIGHT G. N. Lance	91
RESEARCH REPORT A. W. Haslett	104
CORRESPONDENCE	117
SOME BOOKS RECEIVED	123
ABOUT OUR CONTRIBUTORS	127

30 illustrations are between pages 64 and 65

OCEAN-BED PROSPECTING

T. F. GASKELL AND M. N. HILL

SINCE more than two-thirds of the earth's surface is covered by the oceans, it is apparent that for the elucidation of the many problems of the structure of the earth's crust information from the continental areas alone would be inadequate. For example, without knowledge of the basic structural features of both the oceans and continents, it will be impossible to find the depths to the layers at which the rock types are world-encircling, and thus the characteristic differences between the oceans and continents. Likewise in explaining the formation of those topographical features, such as volcanoes and mountain ranges, which are found both at sea and on land, it is essential not only to know their structure but also the common features of the earth's crust from which they are derived. These gaps in our knowledge provide, basically, the reasons for the exploration of the deep ocean by geophysical methods.

The methods used are, in principle, essentially the same as those available for continental exploration. In the practical details of the apparatus and in its application there are, however, considerable differences. There are two reasons for these differences. Firstly, the apparatus has to be designed to operate from a continually moving platform; and, secondly, the area to be explored probably lies at a great depth below the observer.

A further difficulty in deep-sea prospecting lies in our inadequate knowledge of the sea-bed topography. Soundings are widely spaced and numerous features which by continental standards are spectacular still remain to be found. As an example of this, Figure 1 shows the depth profile along a 45-mile line in the Eastern Atlantic basin which had been selected from known soundings, as being flat.

This inadequacy of knowledge results from the lack of interest

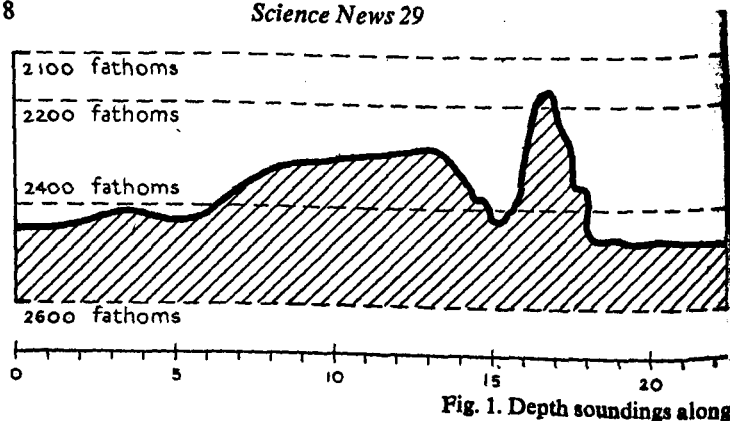
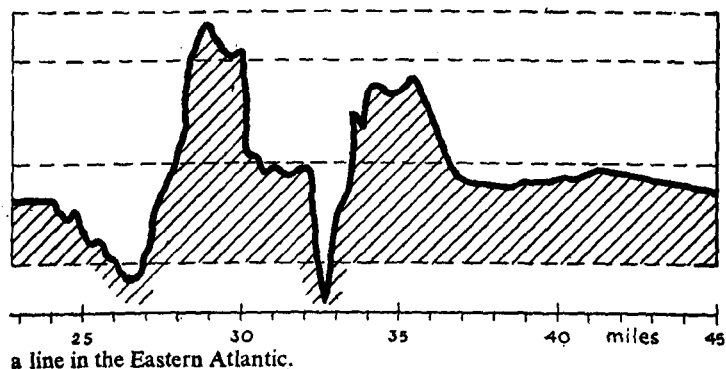


Fig. 1. Depth soundings along

in the deep seas from the point of view of the navigator, and from the difficulties of obtaining deep soundings with a lead. With modern echo-sounders these difficulties have been removed, although the small number of ships fitted for deep-sea echo sounding – there are only the survey and research ships – results in the amount of information obtained being small in comparison with the vast area of the deep ocean. In any geophysical work time must, therefore, be spent in learning about the topography of the area of the experiments. Without this it is impossible to ensure that the area typifies a particular structure.

A second requirement in prospecting in any area is to find the constitution of the surface layers of the sea bed by direct sampling. This is a laborious process and cannot be undertaken at many positions. The labour lies in lowering and hauling the sampling device to and from the sea bed, and the winch required for this work, if a sample of more than the top few inches is to be obtained, is special equipment which is only available in a small number of ships. With the recent development of the Kullenberg 'piston core-sampler', vertical sections through the soft clays and ooze of the sea bed up to 70 feet in length have been obtained by the Swedish Research Ship *Albatross*, although in other ships the handling of coring tubes of this length has proved impracticable and 30-ft cores are more usual. In most of



a line in the Eastern Atlantic.

those parts of the ocean where sediment exists, even the 70-ft cores are relatively only scratches on the surface, and, since the Kullenberg device is, in the best of weather, difficult to handle, smaller corers fitted with coring tubes up to 10 feet in length are useful adjuncts, and by the ease with which they can be handled allow more frequent cores to be obtained. The method which we find most satisfactory for sampling with the small corer is to allow the winch to run freely when the corer is between 50 and 100 feet above the bottom of the sea. This allows enough speed to be gained for the corer to penetrate its full length into the soft materials of the deep sea, and even in shallow water, when rocks are sometimes exposed on the bottom, samples can be obtained. Such rock samples may not be more than a few inches in length, but nevertheless sufficient quantities are usually obtained to allow recognition of the formations to which they belong.

The first geophysical prospecting at sea was the measurement of gravity with the pendulum apparatus developed by Professor Vening Meinesz. The measurement of gravity allows, roughly speaking, the determination of the relative excesses or deficiencies of mass in the layers below the position of the observer. Thus, for example, supposing the rock structure below the sea bed were similar to that under the continents, then between an observer on the surface of the ocean and the centre of the earth

there would be a mass deficiency due to the lighter water replacing the heavier surface rocks of the continent. The value of gravity at sea would therefore be less than would be measured at sea level on the continents. The magnitude of this difference would be such that in the Atlantic the negative anomaly would be about 300 milligal,* while Vening Meinesz by experiment showed the anomaly to be positive and small. This means that the mass deficiency of the sea water must be compensated by the denser rocks which lie deep below the continents coming closer to the surface under the oceans. This conclusion confirms that the surface features of the earth are, in general, *isostatically compensated*, that is floating on the fluid substratum according to the principle of Archimedes.

The gravity measurements were a triumph over extraordinary difficulties caused by the motion of the sea. It was clearly impossible to time the swing of pendulums to an accuracy of a few parts in a million from a surface ship, and Vening Meinesz used a submarine. Even when submerged the motion of the submarine was such that straightforward pendulum measurements similar to those made on land proved impossible, and special apparatus was designed. This apparatus primarily involved the timing of the mean period of two pendulums swinging in the same plane and in antiphase which eliminated to the first order of small quantities the horizontal accelerations, while the timing continued for sufficiently long to eliminate the vertical accelerations which must, unless the mean depth of the submarine changed, average zero. A further requirement in the measurement of gravity at sea is the measurement of the East-to-West velocity of the submarine. This is necessary because of the change in the centrifugal force which is caused by an East-West motion; the magnitude of the effect is such that the velocity must be measured with an accuracy of about $\frac{1}{4}$ knot. In shallow water, where the rate at which wire attached to a weight on the sea bed is streaming out of the submarine can be measured, this velocity correction can

*1 milligal is an acceleration (due to gravity) of 10^{-3} cm./sec.² The value of gravity at sea level is, in round figures, about 10^6 milligal.

be determined with accuracy. In the deep sea, however, this effect may well present the greatest error. A further correction to the observations arises from the comparatively rapid variation in the value of gravity with the submarine's depth, but the determination of depths can be made with great accuracy.

Seismic methods of prospecting probe much deeper into the sea bed than the best of corers, and can provide detailed information about the different layers of rocks instead of the average mass of the geological columns that is given by gravity observations. Sound pulses produced by explosions in the sea are returned from the layers below the ocean bed and picked up by sensitive receivers placed near the sea surface. Two types of experiment may be carried out, seismic *reflexion* shooting and seismic *refraction* shooting. The first type is similar to echo sounding. Whenever in the layers lying below the bed of the sea there are sudden changes in the acoustic impedance (that is in the product of the density of the material and the velocity of sound waves in it), there is a reflecting boundary and sound waves from an explosion in the sea incident on such a boundary will be reflected and can be picked up by hydrophones also in the sea. The disadvantage of this method of prospecting lies in the number of reflecting horizons which are found, and also in the fact that all that can be obtained from reflexion shooting is the time between the explosion and the return of the echo, which in itself cannot show the depth to any reflecting boundary or the change which occurs at that boundary. To determine the depth one must know the velocity of compressional waves in the medium above the boundary; or, conversely, to determine the velocity one must know the depth. Both these are variable over a wide range and neither is, in general, known for the layers lying below the deep sea.

The second type of experiment, seismic *refraction* shooting, allows us to determine both these unknowns, and for this reason its greater complexity justifies its use. With this method the discontinuous changes in velocity which exist between different rock types are determined, as well as the depths at which these changes occur.

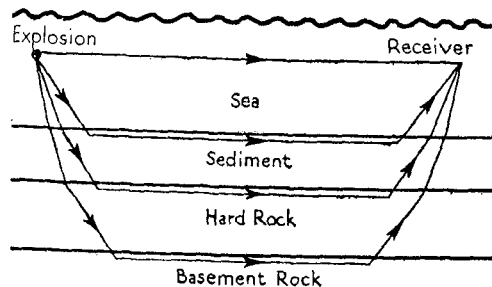


Fig. 2. Sound-ray paths from an explosion to a hydrophone in the sea.

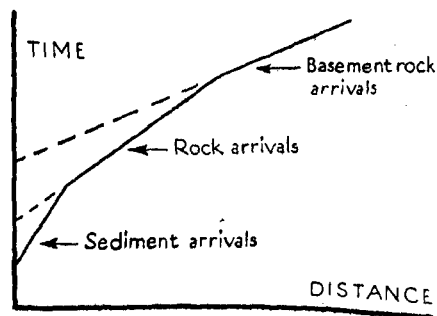


Fig. 3. Idealized time-distance graph. Waves received directly through the water are here ignored (see following discussion).

plotting times of arrival against distances a series of straight lines is obtained, of which the slopes are a measure of the velocity of compressional waves in the horizontal part of the ray path, while the intercept of the straight line on the time axis is a measure of the time that the sound spends on the inclined part and therefore of the depth to the interface.

A pulse of sound arriving by any particular path produces a train of waves through irregularities in the material through which it is passing. The effect of this is that those sounds which arrive by paths other than the shortest-time path are often obscured or their times of arrival are difficult to determine. For

this reason the distance between the shot point and the hydrophone is increased to the extent necessary to ensure that the rays which have travelled to the deepest layers arrive by the shortest time path. This can usually be achieved since, as the range increases, the time taken in reaching any particular layer becomes a smaller proportion of the total time. Hence, because the sound travels with a greater velocity at the greater depths, increasing the range will result in progressively deeper layers providing the shortest time paths. There is one layer, in particular, which never provides the first arrivals. This is the sedimentary layer on the bed of the deep sea. The reason for this is that its thickness is small in comparison with the depth of the sea and the velocity contrast is likewise small. Hence in these experiments the shorter-range shots provide the first arrivals directly through the water, while the longer range ones provide those from layers below the sediments. As a result of this we are confronted with the difficulty of knowing accurately the velocity of the sound in the sediment; in order to determine this we rely upon the uncertain arrivals of sound after those of the shortest time paths and on laboratory measurements on the acoustic properties of deep sea sediments.

In work of this nature on land, accurate surveying is necessary in order to determine the distance between the shot-point and the receiver. At sea, this distance can be measured with extreme accuracy by timing the arrival of the direct sound through the water in which the velocity is known. Another difficulty of the work on land is the boring of the holes down which the charges must be placed. At sea, however, the charge can easily be tamped by firing it at any suitable depth.

Ashore, use is made of moving coil instruments, called geophones, which respond to the velocity with which the ground surface moves after an explosion. At sea, it is clear that similar instruments cannot be placed at the surface but must be laid on the bottom. This is possible in shallow water (that is on the continental shelf, where the depth of water is less than 100 fathoms), but where the depth exceeds a few hundred fathoms it is not generally practicable, although a few experiments have

been undertaken by the Americans in which a geophone and its associated recording equipment have been laid on the bed of the deep sea and recovered by having a float filled with petrol to provide buoyancy. Normal flotation with air-filled containers is impossible because of the high pressures at the sea bed which would cause their collapse.

Geophones as instruments are therefore not generally applicable to deep sea prospecting and we must resort to hydrophones which are positioned some hundred feet below the surface, and which respond to the sound wave pressures. There are two reasons for placing them at this depth: first, so that the surface disturbances shall not affect them; and, secondly, so that reflexions at the surface shall not destructively interfere with the more direct sound pulses. This interference occurs if the hydrophone depth is less than a quarter of a wavelength. In this work the sound frequencies that are used lie between about 15 cycles/sec. and 200 cycles/sec.; the higher frequencies being rapidly attenuated in the solid materials forming the sea bed. The corresponding wave-lengths in water lie between about 300 ft and 25 ft, and we must therefore have the hydrophone at a minimum depth of about 75 ft.

It is also necessary to have hydrophones spaced at various distances from the explosions. The obvious method of achieving this is to have two ships, the first to fire the explosions, and the second for recording the arrivals of the sounds picked up by a hydrophone. This method is used by the Americans and was used on this side of the Atlantic in the first experiments carried out by E. C. Bullard and Gaskell before the war on the continental shelf. The difficulty of obtaining two ships, however, has resulted in a new technique being developed in which hydrophones hang below sono-radio buoys, from which the sounds picked up are transmitted back by radio to receivers in the ship from which the explosions are fired (Figure 4). While this technique was initially developed on account of the expense and difficulty in obtaining two ships for this work, it has subsequently been shown to have other advantages. Thus by the use of more than one of the sono-radio buoys (we usually launch

four simultaneously), several sets of observations can be obtained from a single shot, and thus fewer shots are required for determining the full structure. Another, not so obvious, advantage is that the determination of the dip of the beds underlying the sea is possible.

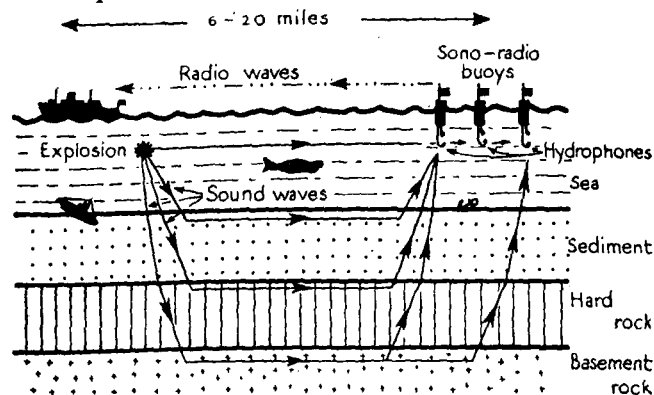


Fig. 4. The single-ship, sono-radio buoy method of seismic refraction shooting.

The method, however, has disadvantages. The first of these lies in the weather restrictions in the use of the buoys. In winds of more than about 15 knots the disturbance caused in the hydrophones by the sideways drift of the buoys causes an excessive background noise, and the difficulties involved in launching and recovering the buoys without damage are considerable. With the two-ship technique the disturbance caused on the hydrophone by a drifting ship is to a large extent eliminated by paying cable out at the rate at which the ship is drifting during the period when the sounds from the explosion are arriving at the ship; this method is not applicable to the buoys. The two-ship technique also has the advantage that it is more rapid since the time wasted in the laying and recovery of the sono-radio buoys is eliminated.

With the buoy method the stabilization of the hydrophone is

the most difficult problem, since changes in hydrophone depth of more than a few millimetres cannot be allowed. Various methods of preventing the vertical motion of the buoy being transmitted to the hydrophone through the cable by which it is supported have been tried. That which has proved most satisfactory is shown in Figure 5. The buoyancy of a float attached to the hydrophone is adjusted so that it can just support the hydrophone and half the submerged weight of the length of cable between the hydrophone and the weight. Thus when the buoy rises and falls in the waves the hydrophone tends to stay fixed. Although this stabilization system is satisfactory for vertical motion of the buoy it does not eliminate the effect of horizontal movements. These do not directly produce spurious signals since the hydrophones are made with two crystals which prevent them from acting as accelerometers. Rather is the disturbance caused by the change in hydrostatic pressure due to the changes in depth of the hydrophone, and by the pressures developed on the sensitive faces of the hydrophone by movement through the water in a direction at right angles to these faces.

A diagrammatic section of one of the buoys is shown in Figure 6. The sounds received on the hydrophone are amplified in a head amplifier on the quartz crystal hydrophone and then in the main amplifier in the buoy. The output of the amplifier is connected with the frequency modulation transmitter which radiates about 10 watts of radio frequency power through the aerial on the flagpole on top of the buoy. Each buoy transmitter has its own receiver in the ship and the sounds are recorded on moving photographic paper on which time-marks are superposed at the rate of 20 per second.

The range of radio reception from the buoys is about 20 miles. This restriction is caused by the necessity of working at a short wavelength in order to have an efficient transmitting aerial, with the result that the limit of reception is approximately the optical range. This horizontal range allows prospecting down to a depth below the sea bed of a few miles. For longer range prospecting with the sono-radio buoys it is necessary to use two ships and the method is shown diagrammatically in Figure 7. This method

has only been used on one set of experiments with R.R.S. *Discovery II* and H.M.S. *Scott* during the summer of 1952.

In the deep sea the spacing between the buoys is varied according to the length of the line along which the shots are to be fired and lies between about one mile and six miles. The first shot is fired between four and eight miles from the nearest buoy and thereafter at similar intervals until the maximum possible range has been reached. For work in shallower water the spacing of both shots and buoys is reduced.

For the deep sea work the explosions are usually produced by depth charges set to fire at several hundred feet. These charges have amounts of explosive in them between 100 lb. and 300 lb. For some of the deep sea and for all the shallow water experiments, however, we use charges which are made up aboard ship and suspended at an appropriate distance below a float. This float is allowed to drift away from the ship and the charge is fired electrically when it is at a safe distance.

RESULTS

The early seismic refraction experiments yielded details of geological structure only in the shallow seas, and the first problem to be tackled by the method was similar on both sides of the Atlantic. The continents are in most parts of the world fringed by an area of shallow sea, which may extend several hundred miles from land as in the English Channel, or may be narrow as along the west coast of America. The change from shallow water of less than 100 fathoms in depth to the normal deep ocean of 2,500 to 3,000 fathoms occurs sharply; this may be because the edge of the permanent continental rocks which do not exist in the deep ocean has been reached, or it may be that the steep slope represents a rubbish tip, where sediment carried down in shallow water is deposited. Measurements by Bullard and Gaskell and by Hill and A. S. Laughton of the thickness of, and the velocity of sound in, the sea bed from dry land to the 100-fathom line in the English Channel showed that a material with velocity ranging from 6,000 to 16,000 feet per second (ft/sec.) covered the hard rocks. The cover material could well be sediment,

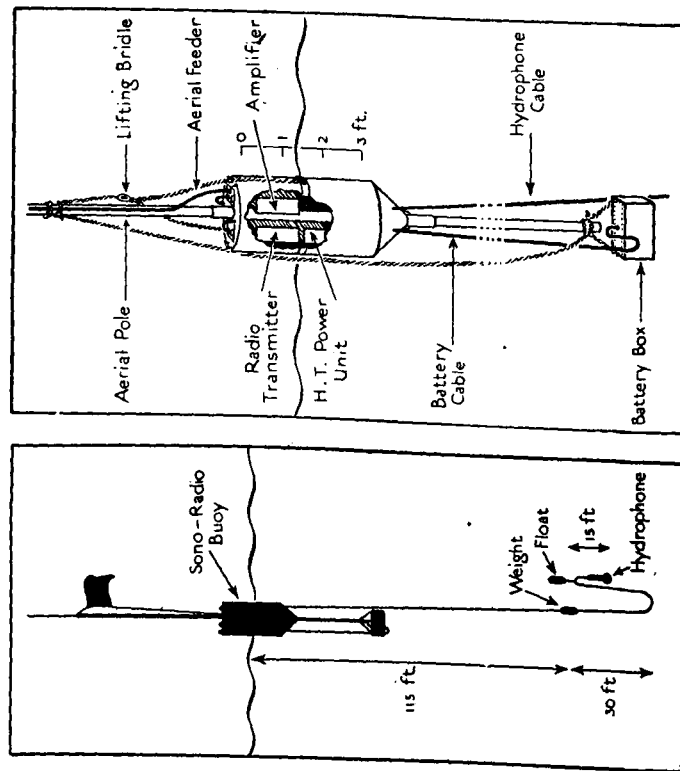


Fig. 5 (left). The stabilization of the hydrophones.

Fig. 6 (right). A sono-radio buoy.

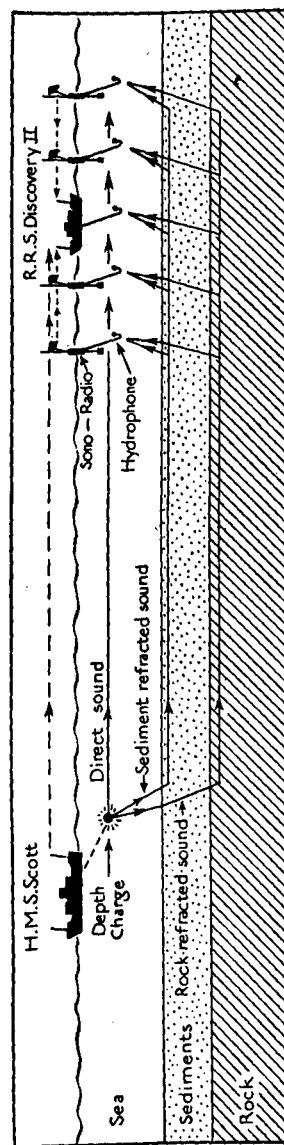


Fig. 7 (above). The two-ship, sono-radio buoy method of seismic refraction shooting.

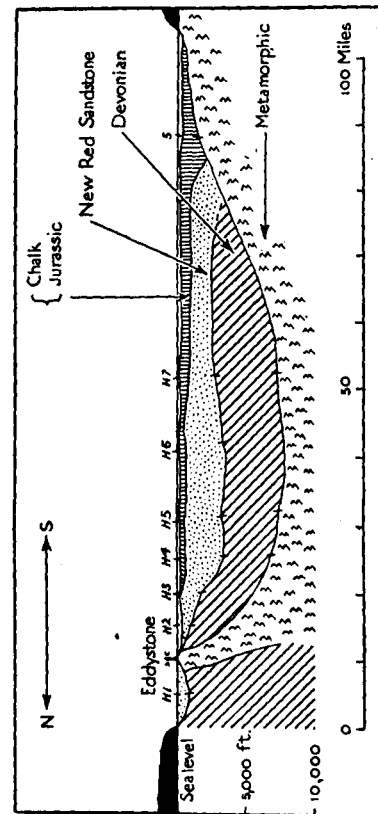


Fig. 8 (below). A geological section through the English Channel southwards from Plymouth.

consolidated by its own weight, and its thickness increasing steadily to about 8,000 ft at the 100-fathom line. A rather similar picture of the structure of the continental shelf was obtained by Professor M. Ewing and others off the North American eastern seaboard. There is some evidence that at least in places the edge of the continental shelf is made of hard rock; samples of sandstone have been dredged from the sea bed off the American shelf, and the character of the topography of the ocean bed near the edge of the shelf is, in places, very rugged. It is not inconceivable that the continental mass of rock joins the deep ocean at the sudden change in depth, and that in places the junction is cliff-like, but that the sediment brought down from the erosion of the land smooths off the land-ocean boundary.

The geology of the English Channel area to the east of the Lizard is not unlike that which would be expected for the normal land surface if it had been eroded to the depth of the sea. Refraction shooting has shown that as in Cheshire or in the south of Hampshire there is a deep basin of New Red Sandstone rocks. In places this is overlain with chalk and earlier formations and is covered with only a thin layer of recent sediments. This layer is in many places so thin that Hill and Professor W. B. R. King obtained geological samples from below it by coring, thus allowing more precise interpretation of the evidence produced by refraction shooting. Figure 8 (on page 21) shows the section through the English Channel southwards from Plymouth.

Deep sea seismic refraction measurements using the sonar radio buoy method have now been made in the North Atlantic and in the Pacific and Indian Oceans. Figure 9 shows the track and seismic stations of the recent cruise of Gaskell and J. C. Swallow in H.M.S. *Challenger*, and Figure 10 those of the 1952 *Discovery II* expedition in the N.E. Atlantic. The *Challenger* did not obtain many measurements in the Atlantic, but the same apparatus and technique was used in the Atlantic, Pacific, and Indian Oceans and the Mediterranean Sea, and a useful continuity of results was achieved. The work of Hill, Laughton, and Swallow in the East Atlantic, of Ewing in the Bermuda area, and of R. W. Raitt in the Pacific is in good agreement with

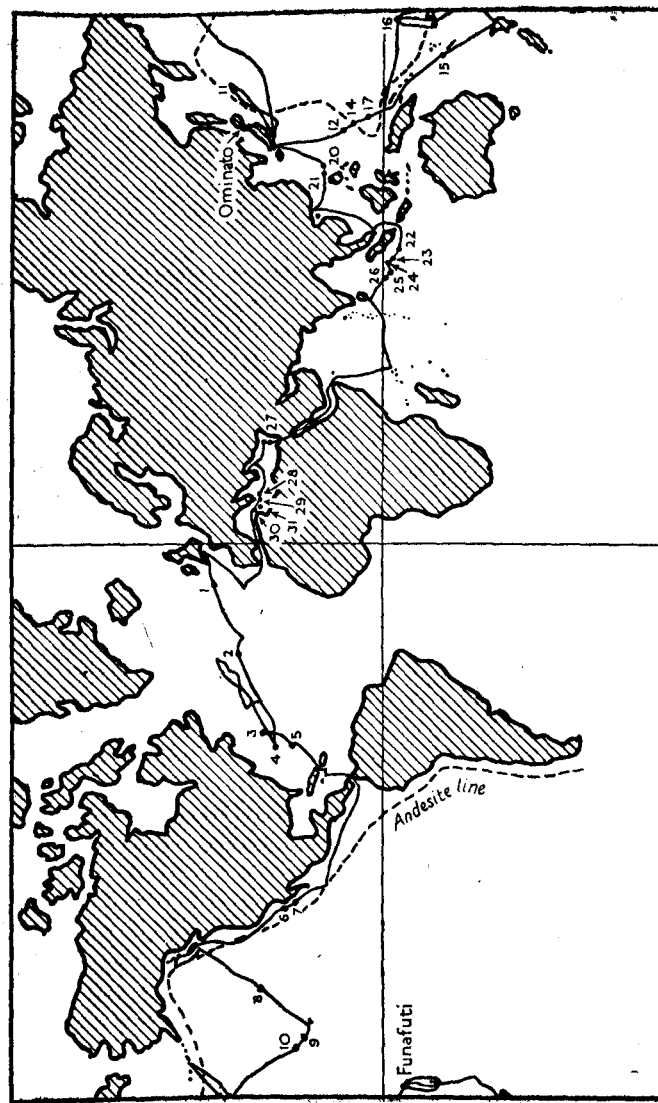


Fig. 9. The track of H.M.S. *Challenger*, 1951-52.

Challenger's observations made by Gaskell and Swallow. Where two ships and large charges are available longer refraction lines than those possible with the single ship sono-radio buoy technique can be undertaken and information about deeper layers in the sea bed can therefore be obtained. However, in most cases the single ship results are adequate to determine the type of geological structure, and the results of the various expeditions can be agreeably combined.

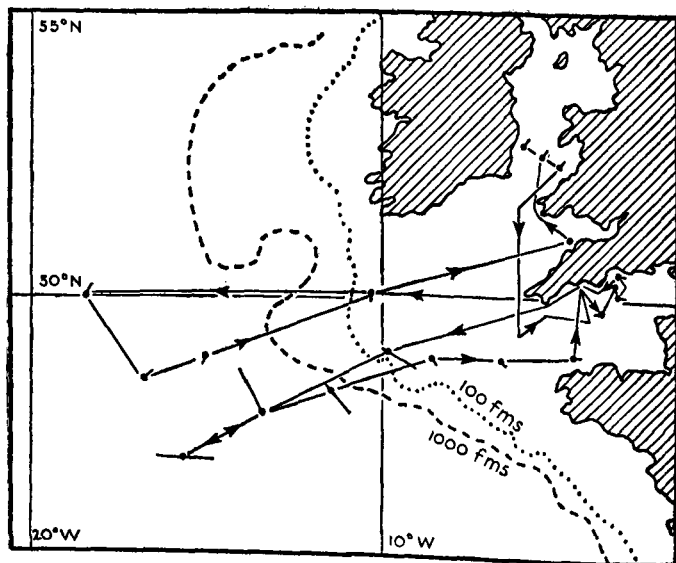


Fig. 10. The track of R.R.S. *Discovery II*. Seismic stations are shown by dots and transverse lines.

There is very little doubt that the present large areas of deep ocean have always been oceans, and have a different geological structure from the continents. Before seismic refraction measurements were made in deep water, the so-called Andesite line had been drawn in the Pacific Ocean to mark the different types of island structure. To seaward of the Andesite line islands

consist of basic volcanic rocks, while on the other side the rocks are more acidic in character and similar to those in the continents themselves. Seismic experiments in the Pacific show that a different type of sea bed exists on the two sides of the Andesite line. To seaward there is what we may call the typical world-wide simple deep ocean structure, in which the ocean floor consists of a few thousand feet of sediments lying on top of a hard layer characterized by a sound velocity of 21,000 ft/sec. which may correspond to basalt. On the continental side of the Andesite line the sound velocity of the basement material is 17,000–19,000 ft/sec., and this is overlain by layers in which the velocity is 10,000–13,000 ft/sec.; similar, in fact, to normal continental structure. The grouping of the seismic results was on the basis of the velocity profiles observed, but this grouping was the same as the geographical distribution given by the Andesite line. A third difference between the deep ocean and the continental type of structure is indicated by Raitt's work in the Pacific, and by Hill and Laughton and Ewing in the Atlantic; at a depth of about 40,000 feet below sea level the basement layer changes to one in which sound travels at 27,000 ft/sec. on the continents. The position of this change is called the Mohorovičić discontinuity. It occurs at about 100,000 feet below the surface in continental areas. This difference between ocean and continent (Figure 11) is perhaps the most fundamental of all, and we believe that below the Mohorovičić discontinuity the rocks, probably peridotitic in character, are world-encircling.

In many parts of the ocean there are large submarine topographical features and there are a large number of islands and island groups standing out in what we have been calling the typical deep ocean. Measurements near such islands show that interposed between the basement rock and the sediment is a material with a sound velocity of 13,000 to 15,000 ft/sec. It is probable that this material is volcanic rock, and a recent laboratory determination of the sound velocity in a sample from the 1950 Mauna Loa lava flow in Hawaii gives 13,500 ft/sec. The velocity could, however, correspond to sedimentary rocks of many sorts; but shales, limestones and the like

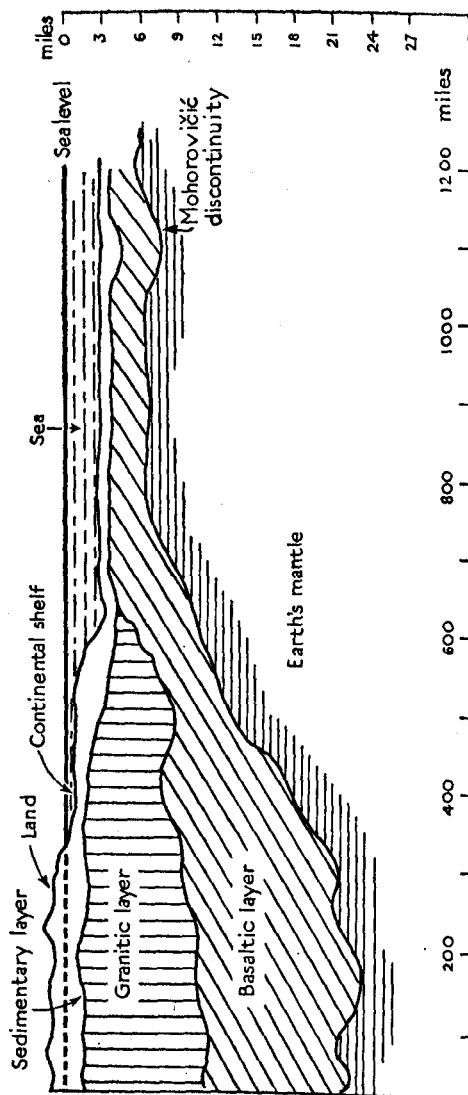


Fig. 11. A section from continent to deep ocean.

must be derived from some source material, and such areas as the Hawaiian chain of islands and the Bermuda area appear to be isolated from land by large expanses of deep ocean having a simple geological structure unlike that of continents.

The volcanic layer is many thousands of feet thick, and such that it penetrates into the normal level of the flat-ocean basement rock. It is therefore probable that the crust forming the sea bed has been depressed, either because material has been removed from underneath or because weight has been placed on top. The sequence of events is, then, that a volcano has erupted through the 21,000-ft/sec. layer forming the bed of the sea, has gradually built itself up and in the course of time has sunk because the void which was left by escaping material must be filled, and because the earth's crust could not stand the strain of the extra weight when sea water was replaced by rock. Measurements at the coral atolls of Funafuti and Nukufetau in the South Seas indicate that these islands are built on old volcanoes, and that there is a core of high-velocity basement material which has not been altered by air or water in the way that the low-velocity volcanic rock which has flowed to form the shoulders of the volcanic pile has been changed. The oceans contain many isolated islands, and as deep ocean surveying produces more details of the sea floor, so more and more isolated 'near' islands or sea mounts become apparent. It is unlikely that these are formed by sedimentation, since they are nowhere near any source of material; yet they often rise 10,000 feet from the flat sea bed, and although sediments of much greater thickness are known in many parts of the world, they are always to be found in areas where an abundance of obvious source material has been available in the past. It is, of course, possible that large features, such as the submarine ridge running nearly at right angles to the Hawaiian chain or the mid-Atlantic ridge, are large folds in the sea bed, because presumably this can be folded in the same fashion as can rocks on continents, but the present known facts do not require any further explanation of deep ocean structure than the simple picture of a permanent ocean bed, consisting of basement rock with a sound velocity of 21,000

ft./sec., studded with volcanic protuberances and their allied growths of coral and debris.

In the south of the Indian Ocean, to the west of the Sunda Straits, the *Challenger* expedition found the simple deep ocean structure consisting of 1,500-ft thickness of sediment lying on basement rock, but to the north of 58° S. there is a considerable thickness of material with the intermediate velocity of 13,000–15,000 ft/sec. interposed between sediment and basement. The Indian Ocean in this area is extraordinarily flat; for hundreds of miles there are no irregularities of more than a few fathoms except for isolated sea-mounts, one large one rising steeply to 7,000 ft above the sea bed, and extending for about 60 miles. This sea-mount consists of 1,500 ft of sediment above a 14,000-ft/sec. material, and the whole area north of 5° S. in the Indian Ocean may be an enormous outpouring of volcanic material. On the other hand, it may be continental in character, and only determination of the depth to the Mohorovičić discontinuity can decide this question.

The sediment thickness in the N.E. Atlantic is greater than in the Pacific, by a ratio of about $2\frac{1}{2}$ to 1. This may be because the Atlantic stations are nearer to the continental source of sediment material, but it may be that there is a more fundamental difference provided by different sedimentation processes; most of the Pacific area is red clay, while in the N.E. Atlantic the sediment is predominantly globigerina ooze. Recent examination of samples from the N.E. Atlantic suggest that some sediment has been derived from the melting of icebergs, and the extra thickness may be related to proximity of glaciers during the Pleistocene period.

In the past, measurement of the present rates of sedimentation of deep ocean oozes has been our most powerful method of estimating the thickness of the sedimentary cover of the rock layers lying deep below the sea bed. There can be little doubt that the approximate rate of deposition is at present about 0.4 inch per 1,000 years. If one assumes that this rate has continued since the ocean basins were first formed about 2,000 million years ago, and allows for the compaction which must exist in the

deeper layers of the sediment, then one would expect a total thickness of about 30,000 feet. This is greatly in excess of the average of about 2,000 feet measured by seismic methods, and we therefore can conclude that the present rate of deposition is greatly in excess of the average taken over geological time. There are several reasons why this may be the case. Firstly, the present relief of the continents is higher than the geologists believe to have been usual in the past. This means that, other things being equal, the rate of erosion will be higher, and the fine clay particles which are deposited in the deep sea will be greater in number. Secondly, a large proportion of the deep sea deposits is organic in origin, consisting of the shells of planktonic animals. The sedimentation rate must depend upon the population of these animals and there are many reasons why this may have been less. Thirdly, it has been found that the small shells are, in many parts of the oceans, dissolved before reaching the bottom and this condition may, in the past, have been more prevalent than today.

The deep oceans and the continents are different in their geological structure, but our knowledge of how the continent joins the ocean is limited. In particular, there are areas such as the Philippine Sea where the water is as deep as in the main part of the Pacific Ocean, and yet from the geology of the islands the area has been shown to consist of continental rocks. Seismic experiments suggest that there may be a gradual transition from continental to deep ocean structure, but long distance refraction results will be needed to be certain of this.

Coral grows in tropical areas in shallow water, and in the Pacific and Indian Oceans there are many coral atolls that are ring-like islands, enclosing a lagoon. Darwin saw many of the Pacific atolls in his celebrated voyage in the *Beagle* in the early 19th century, and concluded that the coral had grown on a sinking volcano; in the early stages the ring of coral had been a reef fringing the volcanic island, and the coral had grown to keep itself above low water, while the volcano gradually disappeared from sight. Murray, on the other hand, decided from evidence in the Solomon Islands that the coral was only a layer of a few

hundred feet that had grown up on clay-like sediments covering submerged sea-mounts; the ring structure was due to the dislike of the coral for the muddy water, so that it only grew on the periphery of the sea-mount. Since reef building corals do not live in depths greater than about 200 ft one would expect no more than this thickness on Murray's hypothesis, while if subsidence occurs a very great thickness may be formed. Evidence from seismic refraction measurements, collected recently by Gaskell and Swallow in the Pacific, shows that the coral is normally several thousand feet thick, so that the typical isolated coral atoll in the deep oceans is formed as Darwin suggested, and as would be expected from the structure of the Pacific Ocean floor mentioned earlier in this article (Figure 12). It is probable that some atolls of the Murray type do exist near continental areas, where thick layers of sedimentary rocks show that large up and down movements can occur, and where bands of coral rock resting on clay are to be found on land.

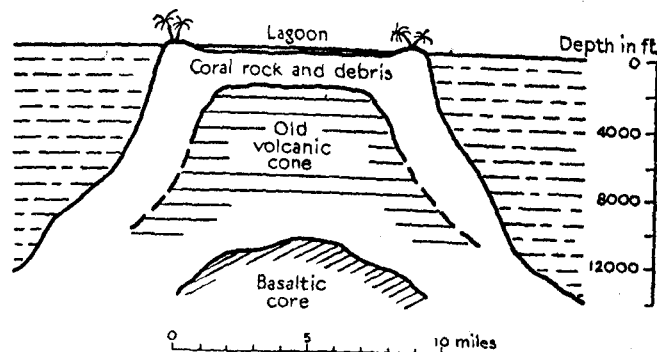


Fig. 12. A section through a coral atoll.

Although seismic refraction measurements have been made in relatively few positions, sparsely scattered over the deep oceans, a reasonable picture of submarine geology can be built up because of its simplicity compared with the complexity of rock layers on the continents. Research continues whenever the

facilities become available for these difficult oceanographic experiments, and will probably yield more information soon on the structure of the fringes of the continents. It is possible that these areas are outgrowths of the continental masses into the oceans; long lines of volcanoes rising out of the deep ocean bed may form traps to hold in the sediment from eroding continents. This material instead of being spread evenly over the ocean floor may thus be concentrated and after hardening by pressure and by heat may be uplifted by earth movements to form extensions to the continents.

Seismic experiments at sea are necessarily carried out from floating, mobile laboratories, and it is by the interest and goodwill of the Hydrographer of the Navy, the National Institute of Oceanography, and the Marine Biological Association of the United Kingdom, that ships are provided so that Great Britain is able to make a large contribution to this world-wide research, and thus carry on the old British tradition of exploring the vast waters of the earth.

THE MOVEMENT OF INDIVIDUAL ATOMS IN METALS

R. S. BARNES

WHEREAS the atoms of a liquid or gas wander about almost freely, the atoms of a *solid* are held more or less rigidly in position by their neighbours, and in a *crystalline* solid, such as a metal, the atoms form a regular pattern or lattice. At low temperatures the atoms vibrate a little about their mean position in the lattice, but as the temperature is raised the average energy of the atoms increases and on the whole they vibrate more violently, until at the melting point the lattice structure breaks up and the material becomes a liquid with comparative freedom of atomic movement. However, even in the solid state some freedom of movement is allowed, for at temperatures below the melting point, although the average energy is not sufficiently large to break up the lattice structure, some few atoms acquire sufficient energy from their neighbours to be able to break away from their lattice positions and to move to new ones, but it is only when the metal is within a few hundred degrees of its melting point that the number of such atoms is appreciable. This process, involving the movement of atoms from one lattice position to another, is known as *diffusion*.

That the atoms of a solid metal diffuse is apparent if we take two pieces of metal, one the radioactive isotope* of the other, and heat them together in contact for a period. In regions where they touched, some of the radioactivity will be found to have

*The radioactive isotope contains atoms which are chemically identical with the atoms of the normal metal, but which are effectively 'labelled' in that they emit radioactive particles which can be detected.

transferred to the other side of the join, showing that radioactive atoms have crossed the join although the metal remained solid throughout the treatment. This process, where an atom diffuses about a lattice composed of atoms chemically identical with itself, is known as *self-diffusion*. Similarly, if two different metals were heated together, they would mix in regions where they touch, but in this case the whole range of alloys of the two metals would be formed in the diffused regions. Such a process involving the simultaneous diffusion of two different types of atoms is known as *interdiffusion*; this is more complex than self-diffusion.

As the diffusion of atoms is the controlling factor in many metallurgical processes it is very important to know at what rate the atoms diffuse under all possible conditions. For example, an alloy when cast is usually inhomogeneous and must be annealed at a high temperature, so that the atoms can diffuse to give it a uniform composition throughout. From the rate of interdiffusion the necessary time for such an anneal at any temperature can be estimated. Also it is frequently necessary to use two metals in contact at high temperatures, and a knowledge of the interdiffusion rates enables a maximum operating temperature to be specified so that interdiffusion of the metals can be avoided and the particularly advantageous properties of each metal retained. Apart from such immediately practical problems, a study of the diffusion process will help our understanding of the behaviour of metals at high temperatures, and make it possible to develop specialist alloys for high-temperature operation; alloys to operate at progressively higher temperatures are constantly being demanded by the designers of gas turbine and rocket engines.

HOW DO THE ATOMS OF A METAL MOVE?

An atom will in general move about the lattice by taking the path of least resistance for which it needs least energy. There are *three* reasonably possible ways in which an atom can change its lattice position. The first process is shown overleaf in Figure 1, where two neighbouring atoms (*a* and *b*) squeeze past one another to exchange positions; this process we shall call the direct

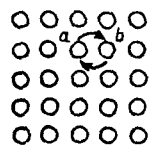


Fig. 1. Direct interchange mechanism.

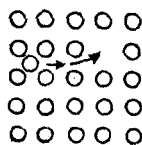


Fig. 2. Interstitial mechanism.

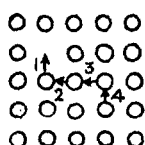


Fig. 3. Vacancy mechanism.

interchange mechanism. Figure 2 represents a process in which a single atom squeezes *between* neighbouring atoms, to wander about until it arrives at a vacant lattice site (or vacancy); there are always a few of these vacancies in any metal lattice. This process is known as the interstitial mechanism, because the atom travels from one interstitial position to another until a vacancy is found. The final process (the vacancy mechanism, see Figure 3) consists of the movement of an atom into a *neighbouring* vacancy; another atom moving into the newly vacated site. In effect the vacancy diffuses through the lattice, although this results from the series of short jumps of many different atoms.

HOW FAST DO ATOMS MOVE?

It is not possible to study the movement of a single atom in a metal, but we can study the effect produced by the diffusion of a large number of atoms. For instance, if we take a thin sheet of a metal and deposit upon its upper surface a very thin layer of its radioactive isotope (see Figure 4) and heat the sheet to a

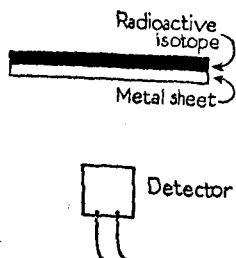


Fig. 4. Arrangement for measuring rates of self-diffusion.

sufficiently high temperature, self-diffusion will occur, and the labelled atoms will move into, and about, the lattice of atoms in the sheet. If the radioactive isotope chosen emits radiations which penetrate only a short distance through the metal, a detector placed below the sheet (Figure 4) will not record until some of the labelled atoms reach the underside of the sheet. The faster the atoms diffuse, the sooner the detector on the other side of the sheet will record.

Such an arrangement could be used to show the way in which the rate of diffusion depends upon the temperature. It is found that the time (t) for the atoms to diffuse through the sheet is an exponential function of the absolute temperature (T)

$$t = t_0 \cdot \exp(Q/T)$$

(where t_0 and Q are constant for a particular metal) so that at low temperatures the time for the atoms to diffuse through the sheet is very long, and it is only when the temperature is within a few hundreds of degrees Celsius (centigrade) of the melting point that the time is reasonably small, and diffusion appreciable.

A series of such experiments all done at a given temperature but on several sheets each of a different thickness would naturally show that the thicker the sheet the longer the time the atoms would take to diffuse through it. In fact, the relationship between the time (t) and the thickness of the sheet (x) is given by the relation

$$t = x^2 / D$$

(where D is constant for a given metal). The value of D for a metal can be used as a measure of the ease of diffusion through it when comparing it with other metals.

The interdiffusion of two *different* metals can be studied in a similar way if some chemical method of detecting the penetration is used.

To give some idea of the rate at which metals interdiffuse; after a day at 1,000° C. (84° C. below the melting point of copper) copper diffuses only 0.002 in. into nickel. However, as in many applications metals are maintained at high temperatures for very long periods, such a rate can be very important, particularly if thin sheets are involved.

MOVEMENT ALONG GRAIN BOUNDARIES

The lattice structure of a metal is never perfect, and, as the imperfections in the arrangement of the atoms produce a looser packing, it is easier for the atoms to diffuse in their vicinity. Although the lattice structure is predominantly regular, it is not normally made up of continuous rows of atoms starting at one surface of the metal and ending in another. The atoms are only aligned in small regions called grains; one grain containing rows of atoms lying at an angle to the rows in the neighbouring grains (see Figure 5). Along the boundary between the two grains (the grain boundary) the atoms cannot be neatly arranged, and we should expect atoms to diffuse more easily along such a grain boundary, than through the more tightly packed regular lattice. This is indeed the case.

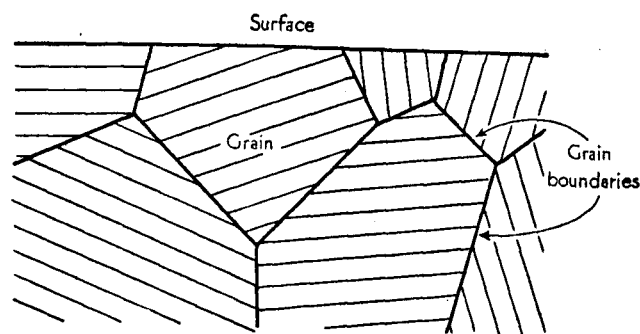


Fig. 5. The grain structure of a metal.

Inset 2 shows a photomicrograph of a diffusion couple (made by welding pieces of copper above and below a thin nickel sheet), which has been chemically treated so that the central nickel has been attacked and darkened. The copper has penetrated deeper into the grain boundaries of the nickel (which are revealed by the chemical treatment as fine black lines) than into the grains themselves, showing that the diffusion has been more rapid along the grain boundaries. Inset 3 is a similar photo-

micrograph, but this time of a surface cut parallel to the welded interface, so that the copper (in this case blackened) is diffusing normal to the paper. It can be seen that the copper has penetrated the grain boundaries, whereas the grains themselves are uncontaminated. Inset 4 and 5 illustrate further examples of grain boundary diffusion (similar to Inset 2); the thin black line (xy), appearing on the nickel-rich side of the interface, marks the extent of the copper penetration. In these two photographs only one interface is shown in each, because the thickness of the nickel sheet is much greater than that shown in Inset 2, where two interfaces appear. The penetration at the grain boundaries is more marked in Inset 4 (where the couple was diffused at a lower temperature) than in Inset 5, where, however, the normal penetration into the grains is greater. It is generally true that, near the melting point, the diffusion along the grain boundaries is but a small contribution to the total diffusion, but at lower temperatures the diffusion along the grain boundaries becomes relatively more important.

Atoms on the surface of a metal diffuse even faster than atoms in the grain boundaries, as is to be expected seeing that they are much freer than atoms in the body of the metal, even those in the grain boundaries.

THE MOVEMENT OF MARKERS

When two metals diffuse one into the other, do they do so at the same rate? The answer to this question can give valuable information on how the atoms diffuse. It can be studied by using a couple, made of the two metals in question (*A* and *B*), such as that shown in Figure 6. The flat join between the two metals is made visible by including markers of some material which does not react with either of the metals *A* or *B*. On heating such a couple to allow the metals to inter-diffuse it is found that the markers move towards one side of the couple (the *A*-side, say) and away from the other (the *B*-side). The higher the temperature the more rapid is this movement, and the total movement at any one temperature is proportional to the square root of the time for which the couple has been annealed. This behaviour

suggests that the movement is the result of the inter-diffusion process, and this is confirmed by making the two halves of the couple of the same metal (or one half of the metal and the other of its radioactive isotope), when only self-diffusion occurs, and the markers do not move.

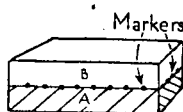


Fig. 6. Diffusion couple with the welded interface marked.

The markers do not move relative to the lattice immediately surrounding them, as markers of any shape or size move at the same rate, and even markers, which dissolve into the lattice, and are so intimately connected with it, move at the same rate. This can only mean that the lattice in this region is itself moving; the markers drifting with it and making the movement apparent.

If more atoms of *A* than of *B* crossed the marked interface, then there would be a steady gain of atoms on the *B*-side of the interface and a steady loss on the other. To accommodate these extra atoms in the predominantly *B* lattice, the lattice would have to expand. Similarly the steady loss of atoms on the *A*-rich side would cause this side to contract. Between them, the expansion on one side and the contraction on the other, they would produce a movement of the markers as is observed, and we must conclude that the movement of the markers is a manifestation of such a preferential flow of *A* atoms over the interface. The marker movement is much larger than could be explained if a direct interchange mechanism was operating. Because then for every *A* atom which crossed the marked interface, a *B* atom would have to cross in the opposite direction, so that no expansion or contraction would result if the sizes of the *A* and *B* atoms were the same. There normally is a difference in size of the two types of atoms, but it is never large enough to produce an effect as large as is observed.

Marker movements occur in all pairs of metals so far tested, e.g. silver/gold, copper/nickel, copper/brass, etc., showing that interdiffusion in all these metals must occur by some mechanism other than direct interchange. Either the vacancy or interstitial mechanism, neither of which necessitates the flow of an equal number of *A* and *B* atoms over the interface, is left as a possible one.

EXPANSION RESULTING FROM DIFFUSION

In addition to the movement of markers it is found that as two metals inter-diffuse a series of small roughly spherical holes appear in the body of the diffusing metal. These holes appear near to the welded interface, but not actually on the interface, and they are always on that side of the interface which is contracting, i.e. the side which on the whole is losing atoms (the *A*-rich side in our example). It is found that the sizes of the holes increase as diffusion proceeds, the increase being more rapid at high temperatures. The two photomicrographs of cross-sections of diffusion couples with their welded interfaces perpendicular to the plane of the photograph, shown in Inset 4 and 5, illustrate the nature of these holes. The welded interface in this case is marked by a series of black dots, which are small inclusions entrapped during welding. The presence of the holes causes the density of the couple to decrease, and the progressive decrease in the density which is observed when two metals inter-diffuse can be entirely accounted for by the formation and growth of these holes. This decrease in density is also proportional to the square-root of the annealing time, which suggests that the hole formation also is a result of the diffusion process. The formation and growth of holes as a result of interdiffusion suggests (as we shall show) that vacant lattice sites are in some way involved in the diffusion process.

INTERPRETATION

The experiments described make it possible to deduce how the atoms behave when the interdiffusion of two metals occurs. The fact that markers move during interdiffusion eliminates the

possibility that a direct interchange mechanism operates, leaving the vacancy and interstitial mechanisms as possible ones. The vacancy mechanism is, however, the more favoured, mainly because calculations show that the energy necessary to form a vacancy is smaller than that to form an interstitial atom. But also the experimental results (in particular the growth of the holes) can be explained more simply and plausibly if a vacancy mechanism operates.

If atoms diffuse by exchanging places with vacancies, the flow of atoms moving in one direction will be accompanied by the flow of an equal number of vacancies in the opposite direction. In the example we have taken (Fig. 7) the preferential diffusion of *A* atoms in the positive direction will be accompanied by a preferential flow of vacancies in the negative direction, so that vacancies will accumulate in the *A*-rich region (where the holes are found).

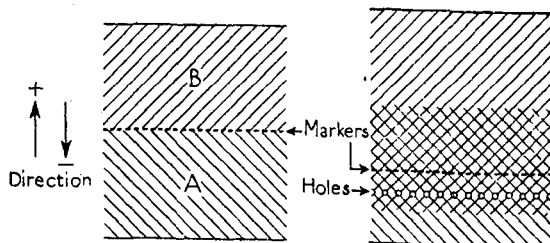


Fig. 7. Couple, left, before diffusion; right, after diffusion.

An obvious question to ask is where does the large number of vacancies come from. The number of vacancies *normally* in the lattice at diffusion temperatures is very much smaller than the number of vacancies which must cross the interface in even a very short time, so that vacancies must be continually generated on the side of the interface from which vacancies are being lost, to make good the loss. These vacancies cannot come from the surface of the couple, because the marker movement is not dependent in any way upon the size or shape of the couple, but they

must be generated in the diffusing region of the couple. A vacancy can be produced most easily, i.e. by using the least amount of energy, at a discontinuity in the lattice where the atoms are less tightly packed. Discontinuities which could act as sources of vacancies occur at (a) grain boundaries, which have already been described, (b) impurity atoms, which occur even in the most pure metals, and (c) edge dislocations. This latter type of imperfection, which occurs in all metals and accounts for many of their mechanical properties, can be regarded (Figure 8) as being caused by the insertion of an extra partial plane of atoms, *l*, between two of the normal lattice planes, *m* and *n*.

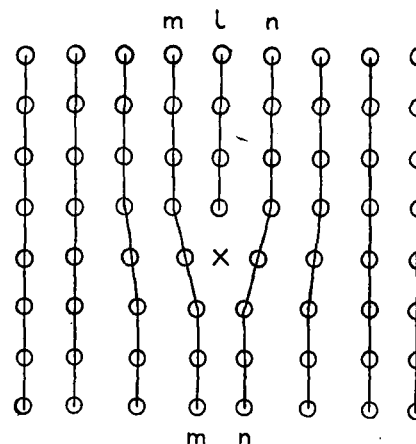


Fig. 8. Edge dislocation.

The separation of the atoms around the position marked X is greater than the normal spacing, because the extra plane of atoms (*l*) forces apart its neighbouring planes (*m* and *n*), making it easier for atoms to move in this region. If an atom in either of the planes *m* or *n*, and adjoining the position marked X, moved into the position X, a vacancy would be left behind which could then diffuse through the lattice. Thus by adding an extra atom to

the extra plane (*l*) of the dislocation, a vacancy is generated with the use of relatively little energy. As further vacancies are generated, the plane of atoms *l* will grow until it reaches the edge of the crystal, when the dislocation will cease to exist, a complete plane of atoms having been inserted. Similarly, vacancies can be annihilated at an edge dislocation, atoms moving *out* of the plane *l* to exchange places with the vacancies, until eventually this plane of atoms *l* is eliminated, and the dislocation ceases to exist.

However, not all of the large number of vacancies which cross the interface are lost at lattice discontinuities, but about one third of them go to enlarge the holes which form on the side of the interface where the vacancies accumulate. For if a vacancy diffusing about the lattice were to exchange positions with an atom on the surface of a hole inside the metal (see Figure 9), the vacancy would be assimilated by the hole, which would be enlarged by the size of the vacancy. Thus any holes in the lattice would be constantly enlarged in regions where vacancies were accumulating, and if a sufficiently large number of vacancies accumulated in a small region, they would coalesce to form holes in much the same way as water molecules coalesce

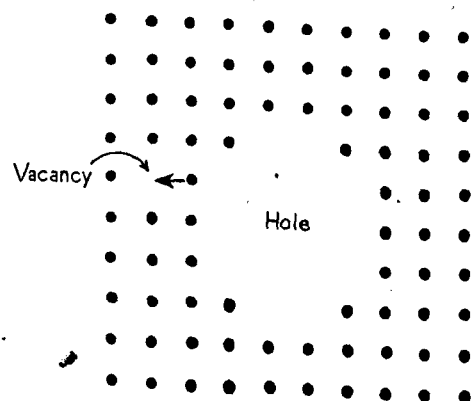


Fig. 9. Enlargement of hole by acquiring a vacancy.

to form water droplets on a misty day. Thus the accumulation of vacancies in large numbers, in the region which is on the whole losing atoms, can explain the formation and growth of holes.

Thus when two metals interdiffuse, vacancies flow across the interface, many more going towards the *A*-rich side than in the opposite direction. This net flow of vacancies results in an accumulation of vacancies on the *A*-rich side and a reduction of the number on the *B*-rich side. The flow of vacancies is maintained by the continual replenishment of vacancies at lattice discontinuities on the *B*-rich side. Some of the vacancies which flow over the interface to the *A*-rich side are annihilated at similar lattice discontinuities, the rest going to enlarge the holes formed in the early stages of diffusion. The reason for this preferential flow is presumably that it is easier for one type of atom (in our case the *A* atoms) to exchange places with a vacancy, so that these atoms diffuse more quickly than the other type.

PRACTICAL APPLICATIONS

Knowledge of the behaviour of diffusing metals can be applied to many metallurgical processes, and two examples will be given. It is common practice to coat the surface of a metal article with a thin layer of a different metal to give the article a better surface finish or higher resistance to corrosion. If such an article is subsequently heated the metals may interdiffuse, and microscopic holes form near to the interface, producing a weak region in the metal. If the holes are sufficiently near together, the surface layer may peel off, exposing the underlying metal. A full understanding of the diffusion behaviour of metals would enable such failures to be avoided. Many small metal components are now made by sintering metal powders, which have previously been pressed into the desired shape. When such a compacted powder mass is heated the individual particles sinter together to form a mass full of holes, which are formed by the particles closing round about them. On continued heating the density increases until the compact is almost solid because of a reduction in the total volume of the holes. However if the powder used consists of a mixture of particles of two different metals, it is found that the

density of the compact at first *decreases*, and only after continued heating does the density start to increase. This initial decrease of density is caused by the formation of the extra holes which result from the interdiffusion of the metals. As the size of the articles produced by such powder methods must in general be controlled, this diffusion effect must be understood.

As a result of these recent investigations of the diffusion phenomenon it has been realized that many early interdiffusion measurements are in error because the movement of the interface and the formation of holes has been ignored. It will be necessary to repeat them if accurate results are to be obtained.

These experiments on the diffusion of metals have given valuable information on the detailed atomic processes occurring in metals at high temperatures. In particular they have shown that vacancies play an important role.

METHODS IN PSYCHOLOGY

F. H. GEORGE

THE psychologist who would present his subject to non-psychologists is confronted with several difficulties. One of them is that every human being is, in some degree, a psychologist: he may be open to the criticism that he studies his own behaviour by introspection and that of others by intermittent observation; but, in his own eyes, and in his own way, he is still something of a psychologist. He expects to be told more about himself than he knows already, or thinks he knows. But the real extent of his knowledge is suggested by the eagerness with which isolated concepts, from sex-motivation to intelligence quotients, are taken up, puffed up, and what may be limited but useful notions become ridiculed, because of the disproportionate importance attached to them. Other difficulties are inherent in psychology as a subject. Ideally, the psychologist should have access to a complete physiology, extending at the least to a full knowledge of all 'physical' changes in the central nervous system, and affecting it. There is as yet no such completeness: there is a gap between the animal-as-a-whole and particular internal processes, not all of which can be directly studied. Should psychology then wait? That is not the normal custom in science. Furthermore, in the practical world it is assumed, in part correctly, that applied psychologists can help the school-teachers with their knowledge of learning, contribute to vocational guidance and the improvement of relations between groups, and even, as psychiatrists, aid in medical treatment. These applications depend on general psychology; the part played in them by special techniques is relatively small, since the principal difficulties lie in interpretation, for which we lack the necessary theoretical model. The psychologist is compelled, therefore, to give practical guidance, though the foundations of his own science are still under discus-

sion. This was for a long time the position also of, for example, the metallurgist. Another difficulty is suggested by some of the definitions given of psychology: terms such as 'conscious life' and 'mind' are a reminder that psychology abuts also on philosophy. Again, this is true of other sciences, and indeed of science as a whole. But in psychology there have been theoretical psychologists who have been philosophical in a sense that the theoretical physicist is not, turning repeatedly as he does to his laboratory colleagues. It is largely as the result of a recent revolt against arm-chair psychologists that the term 'experimental psychology' has come to cover both theory and experiment; its task, as in other branches of science, being to bring about their closer integration. There is also one further difficulty, which is sufficiently suggested by the fact that the Penguin *Dictionary of Psychology* is a fatter volume than its companion *Dictionary of Biology*. In this article I hope, none the less, to clarify the present position, and to make psychology understandable as a science to other scientists.

Such an explanation is the more desirable because the compartments of science are essentially man-made. The subject matter of all science is essentially the same: systematic observation and systematic presentation of the observations in communicable form. Scientists naturally differ in the levels of abstraction at which they approach knowledge; they work on various dimensions and with various techniques. The division of science into compartments, however, is a convenience; it is indeed a necessity as a division of labour, as are all categorizations and classifications. They are also quite arbitrary; and, although such divisions are necessary for discussion and to ease the burden of administration, it is their arbitrariness that has, unhappily, been most often forgotten, with the result that they are taken far too literally. These embarrassments have contributed toward narrow specialization as well as the attendant evils of lack of co-ordination, co-operation, integration, etc. It is in this rather unfortunate, rather distorted background that I would like to consider the experimental psychologist's role in science.

EXPERIMENTAL PSYCHOLOGY

We can trace in experimental psychology many large and prominent threads in the rather complex skein of the subject. Firstly, I shall consider the large field of comparative psychology and its overlap with zoology, physiology, biochemistry, endocrinology, and genetics. There are, in comparative psychology, four principal methods of using animals: (1) For the study of the psychology of particular animals: dog, cat, etc. (2) The comparative method in the sense of the zoologist which involves the comparison of various forms. (3) A type of phylogenetic usage which is designed to trace the evolution of psychological processes from 'lower' to 'higher' organisms, and (4) the use of the animal as a research tool; an instrument for research in problems of behaviour. The last of these methods has been by far the most important for the psychologist. It is, indeed, vital to the work of the psychologist in animal behaviour to realize the basic fact that psychologists have generally used the animal as an experimental tool and not in a truly comparative way. This must be regarded as a criticism of animal psychology to date and has diminished to some extent the value of its findings. The development of a truly comparative method seems essential to the future of psychology.

For experimental psychology the animal represents a means of observing a simplified combination of behavioural patterns. The use of relatively simple organisms has, it is believed, the advantage of reducing the number of variables under consideration and allowing a control of the 'external environmental' variables that could never be possible, say, with humans. This leaves difficulties in dealing with 'internal variables' and the general lack of control that one can exert over them. What is worse, even in physiological psychology, where the organism can be placed under some kind of internal control, the behavioural patterns are to some extent changed by the actual introduction of the control.

The above difficulties seem to be quite general in experiments on behaviour, as they are in all science in some dimension. The

well-known Heisenberg principle of 'Indeterminacy', for example, deals in physics with the fact that carrying out an observation interferes to some extent with what is observed. In the above statement we are adhering to a scientific model of a 'realistic' kind (assuming a 'real' external world), and this is not regarded as inconsistent with an insistence on operational definitions, i.e. definitions which merely refer to the actual operations performed by the experimenter, for certain models on certain dimensions; this applies particularly in limiting or borderline cases. Thus, apart from many debatable questions, the study of behaviour deals with problems similar to those in the rest of science and has similar difficulties; the differences being of degree and dimension. In particular, the psychologist is dealing with a large number of relatively uncontrollable variables whose internal relations appear to be by no means linear.

The above statements on theory-construction refer to the possible division of opinion between 'operationism' or 'positivism' on one hand, and philosophical realism on the other, and they seem central to all scientific theory building. We shall consider later this vital subject.

In these circumstances the psychologist is doing just what one might expect; he attempts to control as many variables as possible, and with the aid of simplified animal experiments and statistical techniques he aims to produce correlations between environmental conditions and behavioural patterns, neural structure, biochemical events, etc., which will lead to inductions and generalized laws of behaviour. It is true that in this sentence the phrase 'to produce correlations between behavioural patterns and neural structure, etc.' is not always regarded as necessary. Much of behaviour theory is said to be 'molar', i.e. independent of the underlying 'molecular' complexes of neurology, endocrinology, biochemistry, etc. It must not be thought that psychologists really believe it to be independent, but some do believe that theories of behaviour can be produced entirely in organism-as-a-whole terms and without immediate dependence on the 'molecular level' sciences. Let us consider a typical animal experiment: a maze situation wherein one may observe a rat run

a maze under more or less complicated circumstances. A 'reward' is introduced into the goal-box, and the performance of the rat will vary from time to time and according to the nature and amount of the reward as well as many other factors which are not necessarily observable directly. Thus we have the well-known situation in which we appear to be able to observe most of the external stimuli conditions and the responses-as-a-whole, although further variables, conveniently assumed to influence the resultant behaviour, are unobservables in the organism-as-a-whole, being largely internal to the organism. This viewpoint, which has often been followed, seems to imply a 'realistic' model. Thus from the organism-as-a-whole viewpoint there exists the need for 'hypothetical constructs', such as 'electron' in physics, or 'ego' in psycho-analysis, to bridge the gap in our observables. The special process of reducing these constructs to the status of 'intervening variable' and operational definition, or the replacement of the constructs by an increased range of observation, appear to be the alternatives at our disposal. The last possibility will depend on increasing our range of observability in the organism-as-a-whole technique (the electroencephalogram, etc.), or on supplementing our analysis of a 'molar' kind by analysis of a 'molecular' kind. This point will be reconsidered later.

The theorist of 'molar' behaviour – the psychological psychologist – is not dependent on the data of the analytical or 'molecular' sciences such as neurology, but he couches his theory in terms of behaviour-as-a-whole. He is rather in the position of owning an unopenable box which performs a great variety of tricks; and, since the box is not openable, he makes inductions about the internal mechanism by observations on the resultant performance only. The best examples extant of 'molar' theories of behaviour are those of C. L. Hull and E. C. Tolman. Both use the method of deduction from limited assumptions, and both attempt to fill the gap in our chain of observables, by the use of logical or hypothetical constructs; a system of verbal labels that may be taken to cover complex organic events, internal to the organism. Hull, proceeding on a highly 'mechanistic' pro-

gramme, fulfils the ambition of the behaviourist J. B. Watson and, taking literally the measurements obtainable from the conditional reflex experiments of both the Pavlovian and the neo-Pavlovian type, outlines a stimulus-response type of theory. This is sometimes compared to a slot-machine in which the stimulus, the coin, goes in, and the response, the bar of chocolate or what-have-you, comes out. Tolman goes a long way along the 'behaviourist' path with Hull, but attempts to include in his theory the complex findings of Gestalt psychology and the work of this school on the situation-as-a-whole, mostly on visual perception. He is thus forced to place more emphasis on the part played by the organism than the passive role assigned to it by Hull. The organism in Tolman's theory has sometimes been compared to a (map-reading) control room, which can modify its behaviour without necessarily going through the overt response. Tolman makes more allowance for the organism's internal solution of problems instead of treating its initial responses as random.

Returning to the general situation in experimental psychology we find, then, a theoretical structure manufactured in symbolic (linguistic, aided by metaphorical models) terms, and standing in much the same relation to experiment as does mathematical physics to experimental physics. The two most recent serious developments have been the gradual extension of experimental evidence from 'molar' to 'molecular' levels (more and more attention is gradually being paid to generalizations of a neurological kind), and an increase in the number of attempts to develop some form of mathematical language for behaviour theory.

D. O. Hebb's recent monograph on behaviour theory was an attempt to build a bridge between physiology and psychology. This work may be taken to be one of the first steps in the integration towards a sort of 'generalized human biology'; part of the approaching unification in science.

The mathematician's work is on a higher level of abstraction and more remote from psychology as such, but mathematical models on certain levels of approximation are already available.

Cybernetics, mathematical biophysics, and the quantification of Hull's behaviour theory, may be cited as examples of the mathematician at work on various levels. Also there are the many attempts to state psychological theories precisely in the terms of mathematical logic. These, however, are on a somewhat different plane, though none the less interesting as experiments in theory.

There exist many examples of these mathematical models in behavioural theory, and they are usually of relatively simple form. They generally express relationships between the principal logical constructs of the system. An example is the starting point made by Ivan D. London for a rather limited investigation into what he calls a class of forgetting curves. He considers the initial state to be defined by the number of originally activated units, these constituting what he calls the 'engram complex'. This is the position at time $t=0$, after which no further learning takes place. He assumes further that the rate of 'engram-inactivation' (loss of 'retention') at time t is proportional to the number of units still effective at that time. This is equivalent to assuming that any unit is as likely to become inactivated as any other, and leads to an exponential form of 'forgetting curve'

$$n = ce^{-at}$$

in which c and a are constants. This simple example must suffice here as an indication of *one* of the possible uses of mathematics in behaviour theory.

RELATIONS WITH PHILOSOPHY

It should be apparent to psychologists, although it has not always been so, that theory and experiment are inseparable in science and that the existence of constructed theories implies deductions, inductions, assumptions and all the paraphernalia of methodology which constitute the philosophy of science. Thus it is that the problems of the philosophy of science are often directly illustrated in psychology; while in physics, apparently little related to psychology, the philosophical problems of quantum mechanics make an interesting comparison with those of psychological theory: a point that has so far aroused little comment.

It is not intended to deal explicitly in this article with the problems of the philosophical psychologist. To do so adequately would demand a historical approach, and only a few points of relatively recent interest can be mentioned here. The well-worn discussion between 'mechanists' and 'idealists' still edges its way through, although the scientist's problem has been made clearer by the notion of model-making. Understanding has been helped especially by the analysis of linguistic 'models', defined as any symbolic structure which uses language – everyday, mathematical, or by way of charts, models, etc. – and from which predictions can be made independently of direct 'empirical observations'. It is irrelevant whether the directives accepted in model-building are 'teleological' (purposive) or 'behaviouristic' (non-purposive). The models are justified, if at all, by the functional relationships to which they lead. Thus the postulational theory of Hull and that of Tolman (still under construction) are independent of the particular prescriptions used in theory-building; and terms involving 'purpose' can be regarded as desirable, or not desirable, as interpretations of these relationships according to individual predilections; such words have principally a heuristic value. Words such as 'purpose' are therefore scientifically meaningless, unless an operational definition of them is supplied, as have Wiener and his associates in the field of cybernetics. They can then be properly used; but thus, and only thus, can they be shown to be necessary to the theory, and not merely metaphysical encumbrances.

The older problem of mechanistic versus purposive directives being largely verbal, the more interesting present discussion is between 'operationism' and 'realism'. I believe that this, too, can be resolved if we regard P. W. Bridgman's operational methods as limiting tests applied to an otherwise 'realistic' model. Herbert Dingle has shown the necessity of the operational method, and its far-reaching successes in relativity theory; it is, I believe, a technique which is equally vital to psychology. It is sufficient to say here that there is no evidence that one model will be universally useful. Use can be found for 'teleological', 'mechanistic', and many other kinds of models, and they should not necessarily

be regarded as mutually exclusive; our problem is to find a sufficiently predictive model for the work in hand, and dependent on the level of abstraction, and the nature of the predictions, desired.

Apart from the search for 'hypothetical constructs' in neurology, there have been interesting developments in 'molar' theory. Three nodal points have arisen which have been regarded as fruitful – 'digging points' in the continuum. They have become known as the 'theory of learning', the 'theory of perception' (these are dug at both 'molar' and 'molecular' levels), and the 'theory of motivation'. These nodal points are obviously not independent, and their separation represents to some extent the historical accidents of development as well as the influence of philosophical directives: they are now in process of integration. It is hoped that this unification will lead to the further development of 'General system theory' as developed by L. von Bertalanffy and others. Indeed this very important approach to psychology has been tapped to some extent already in D. Krech's recent development of a hypothetical construct, which he calls a 'dynamic system'.

The experimental work of psychologists has largely been carried on in this theoretical background, and the theoretical methods used, although varying in detail and technique, have been of the same general form.

Statistical mathematics is a tool to the psychologist, not part of the theoretical edifice. It is used as a weapon for showing relationships between variables either in 'experimental' or in 'uncontrolled' conditions. Standard techniques such as product-moment correlation and analysis of variance are used. Statistical techniques are used often in 'molar' experiments. There have been many thousands of such experiments, and a simple example may be considered. G. C. Grindley, in an early experiment, hoped to show differences between 'rewarded' and 'unrewarded' maze-running in chickens. Omitting details of controls, he used two groups of chickens, one of which was fed grains of rice at the end of the maze-run, while the other could see the rice but could not eat it. The difference between their

later performances was shown to be statistically significant by using a measure of their time through the maze and the statistical method of comparison of means.

The above very simple experiment is 'molar' in the sense that it refers to the organism-as-a-whole. As a 'molecular' example the classical experiments of K. S. Lashley might be given. Very briefly, he took a number of groups of rats, carefully controlled genetically, and taught them some activity, such as simple maze-running. He then performed surgical operations on the rats which involved destruction of various cortical areas. The rats, when fully recovered from the effects of operation, were re-introduced into the maze situation, and from the resulting performances he was able to show that maze-running was possible in the rat irrespective of the particular area of the cortex destroyed, and that performance depended only on the *amount* of cortex left intact. The above two briefly outlined experiments are given merely to illustrate method and many possible criticisms and difficulties are not being considered.

PERFORMANCE AND LEARNING

One special set of methods with a long historical background, and of special interest, goes under the name of the psychophysical method. The names of Fechner, Weber, and other early experimental psychologists are associated with these methods. The well-known Weber-Fechner law is usually stated in the form:

$$R = K \log S/S_0$$

where S is the stimulus, R is the sensation, and S_0 is the smallest stimulus which leads to sensation. This is one rather crude generalization that aims at connecting stimulus conditions and 'sensations'. It has been shown to have a limited application in its original form, but Koenig and Brodhun, more recently, have generalized the theory for the human eye to cover just perceptible visual differences over a greater range of intensities. The psychophysical methods, which generally aim to connect stimulus and response, or stimulus and sensation, have various particular techniques variously labelled: (1) method of constant

stimuli, (2) method of single stimuli, (3) method of limits, (4) method of adjustment, etc. They have in common one inherent limitation; all of them involve, in effect, the identification of the measure of the stimulus with the measure of the response, and these, in general, cannot be the same. At the best, in some circumstances, one will give a certain order of approximation to the other, and the ensuing adequacy for prediction will depend on the dimensions of the particular problem. There has been some success with these methods in dealing with visual and auditory problems, and perhaps also to some extent in 'quality control'.

MORE ABOUT THEORY CONSTRUCTION

I would now like to return to the problem of theory-construction in behaviour theory. The problem of 'learning' appears to be central, and can be said to be largely reducible to the problem of 'reinforcement'. The problem in its simplest terms is to state the conditions that lead to certain behavioural processes being 'stamped in', in the organism's repertoire of possible behaviour patterns, and to others being omitted. Hull and Tolman have approached this problem from their two different philosophical standpoints, and yet in the most recent statements of their theories there is surprisingly little difference between them. They both appear to use mixed 'realist' and 'operational' methods and both are forced to make great use of 'hypothetical constructs'. The present position appears to be that the interpretation of hypothetical constructs in one system is different from the interpretation in the other, and no more. As has been briefly remarked, for Tolman the unobservable gap between stimulus and response has been filled by constructs which are inferred to be more 'central' than those used by Hull. For Tolman the emphasis is laid more on the organism's control of his environment, while for Hull the emphasis is on the environmental control of the organism. These differences of emphasis, and that is all they are, are interesting traces of their rather remote ancestry in extreme 'mechanistic' and 'teleological' viewpoints. It appears to be the influence of operationism on these

theorists which has had the increasing effect of replacing inferences, which earlier were mostly derived from *a priori* theory-building directives, by statements of functional relationships between the 'hypothetical constructs' adopted.

The differences in choice of constructs in these two most representative of contemporary theoretical systems are of extreme interest to specialists; but for a wider audience, and in the present state of biological knowledge, there is probably more interest in their methodological approach than in the predictive powers of their theories.

Perhaps the most interesting aspect of the Hull-Tolman controversy is the way they have moved asymptotically towards each other and a theory of common agreement; both would now explicitly separate 'learning' from 'performance'. It is still, one feels, the atavistic shreds of implications from philosophical directives that stand between them. As Lorenz has pointed out, it is not possible to 'explain' the behaviour of an engine by studying only its piston, but neither is it possible to understand an engine without dismantling it at all. Both processes are necessary for a complete picture, and in their enthusiasm the 'mechanists' and 'teleologists' have forced each other off either side of the sensible road. The use of operationism coupled with a set of realistic many-dimensional models is, in my opinion, an answer to a great many of the troubles; as also is an awareness of the linguistic pitfalls. The Hull and Tolman theories are indistinguishable operationally, which suggests that their ultimate unification depends on increasing the number of behavioural observables. This in turn suggests an attack at all levels by experimental psychologists of all kinds, with physiologists, biochemists, geneticists, endocrinologists, mathematicians, etc., working together with a common plan and an idea of each other's mode of attack. There is need also for a critical and unifying influence by philosophers of science. All of this, and no less, will give us the behavioural answers that we want. One step towards this end will be the philosophical education of scientists; the other, presumably, will be the further education of philosophers in the problems of science. But two facts stand out from current psycho-

logical discussion. One is the need for fitting psychology into a more balanced biological background. For surely a psychology independent of physiology is quite unthinkable in practice; and this, if so, implies a great increase in investigations at 'molecular' levels. The second fact is the need to realize that all science is wholly linked up with theory-building, with the need for 'philosophical' vigilance as a corollary.

In such an article as this there is an obligation on the writer to mention the psychologist's skeleton-in-the-cupboard: 'introspection'. To avoid acrimony, I shall make only a few generalizations. Most of existing psychological theory avoids the problem, since such theory represents generalizations from animal experiments. However, introspection seems best viewed as an easy step from 'consciousness' and without any special mystery, at least, for the psychologist. Introspective reports are usually used in human experimentation and clinical work. However, they are, in a sense, merely an extra experimental control, and certainly do not feature in recognizable form in the theoretical edifice. Although this is clearly an inadequate statement of a complicated problem, 'introspection' does not appear vital at the moment in the building of psychological theory.

One further word of limitation should be added. The integration and unification of science will not be sufficient alone to bring about the whole of the generalized biology that is looked for. It is absolutely necessary, as already mentioned, that our range of observability be increased. To take a simple example: it is almost certain that a detailed understanding of the endocrine glands, as well as special methods for measuring their secretion in the living organism, is essential to our behavioural model. The facts are at the present time that, with the exception of the radioactive tracer method of estimating thyroid output, and some methods of measuring gonadal secretions, there appears to be no direct method at all of estimating endocrine output. The pituitary and adrenal secretions, for example, are not ascertainable and, in general, our understanding of the endocrine model is fragmentary. This is a problem in the main for endocrinologists; but as well as the necessary development of such

dynamic techniques, the development of behavioural theory also will benefit by a closer working together of the various branches of science. In fact, it is probable that integration will actually facilitate the discovery of the necessary specific techniques; certainly it will change their values, and possibly may make some unnecessary.

CONCLUSION

This brief note on psychological methods can do no more than outline what has become a vast, unwieldy, and fast-changing subject. More recently, as in physics, many of the changes have been in interpretation of our knowledge, rather than in experimental technique, although both are necessary. The operational method has done a lot to reduce the amount of nonsense that once passed for psychology or behavioural theory, and I think, above all, it is clear that we have to produce a race of scientific 'all-rounders' as well as the tethered specialist. To this end it is first important that each and every scientist, using, as he does, his own special techniques and theories, should realize in his own branch of work that his techniques are limited and his theories fragmentary; and that they will, at some stage, become integrable into more general theory-constructions and scientific models. The effect of this must, in general, be to change the dimensions of each branch, modify the existing fragmentary theories, and often wholly alter the viewpoint of scientists working in what was some small part of science-as-a-whole.

Psychology has often suffered both from bad publicity and from the fact that it lends itself better than do most subjects to the propagation of nonsense. But out of this uneasy background is springing a healthy, rigorous, and methodologically sound infant which, however, will need to be handled with care for some time to come.

FURTHER READING

General introductions:

- Murphy, G. *An Introduction to Psychology*. New York, Harper, 1951.
Zangwill, O. L. *An Introduction to Modern Psychology*. London, Methuen, 1950.

Learning theory:

- Hilgard, E. R. *Theories of Learning*. New York, D. Appleton-Century, 1948. (Includes accounts of both Hull's and Tolman's theories.)

Philosophical aspects:

- Pratt, C. C. *The Logic of Modern Psychology*. London, Macmillan, 1939. (Out of print.)
Wilkie, J. S. *The Science of Mind and Brain*. London, Hutchinson's University Library, 1953.

Experimental (advanced):

- Munn, E. L. *Handbook of Psychological Research on the Rat*. Boston, Houghton Mifflin Co., 1950.

Physiological psychology (advanced):

- Hebb, D. O. *Organization of Behaviour*. New York, John Wiley & Sons, 1949.

CHEESE BIOLOGY

JAMES A. BARNETT

MANY people wonder how cheeses as different as Cheddar, Stilton, Gruyère, and Camembert, can all be made from cow's milk by processes which are more remarkable for their similarities than their differences. These processes consist in clotting the milk, draining the clotted milk of whey and decomposing the curds which are left. The milk proteins are clotted, either by the action of lactic acid produced by bacteria growing in the milk, or by the addition of rennet. The clot tends to contract and express most of the water and water-soluble constituents of the milk. Finally the relatively dry curds are ripened; this part of the manufacture taking several months for some kinds of cheese.

It is in ripening that a cheese develops its characteristic flavour and texture; and ripening depends mainly on biological activity, principally of bacteria and fungi. Much of the skill of a cheese-maker is his ability to produce curd in which these organisms grow appropriately for the kind of cheese being made. Once ripening has started, the cheese-maker has no means of controlling it, apart from adjusting temperature and humidity. To some extent ripening must be self-regulating: not only must conditions be right for desirable activity, but undesirable changes must be prevented, either by excluding harmful organisms altogether or by making the cheese so that it is unsuitable for their growth.

The question of what gives each cheese its distinguishing characteristics is not simply a matter for idle curiosity on the part of those who enjoy eating a number of varieties, it has been a serious practical problem confronting people who have established cheese industries over the last 50 years. Cheese was made in prehistoric times, and the craft has developed great heights of skill in the creation of such cheeses as English Stilton, Swiss

Emmentaler and French Roquefort. Factory cheese-making has grown up only in recent years, meeting the needs of large town communities, and in countries where there are no cheese-making traditions, as in Denmark, the U.S.A., and New Zealand. This industrial revolution has provided money for research, the biologists being concerned largely with finding out exactly what organisms are necessary for each kind of cheese, what these organisms do, under what conditions they act best and what interferes with them.

HOW CHEESE IS MADE

Lactic cheese

In warm weather people often hang up a muslin bag of soured milk and then eat the resulting 'lactic' cheese. This very simple way of making cheese depends on the presence of certain milk-souring bacteria. The bacteria break down some of the milk sugar, lactose, to lactic acid, and the lactic acid coagulates the main milk protein, casein. Whey runs off from the curds which are retained in the muslin bag. A more sophisticated way of making lactic cheese is to add to the milk special cultures of bacteria, carefully selected for the purpose; these are called 'cheese starters', because they start the souring. Starters, in the form of sour milk containing only the desired bacteria, are used in making commercial lactic cheeses.

Cheddar

Although lactic cheese is made by bacterial action alone leading to the coagulation of milk protein, milk may also be clotted by adding rennet, a commercial preparation of enzymes from the stomach of the calf. This is how junket is made. If the junket is hung in a muslin bag, whey runs out and the curds may be used as the basis of other sorts of cheese. A few cheeses, such as Coulommier, are made from milk clotted by rennet alone. But both starter and rennet are used in making Cheddar, and in fact for most kinds of cheese. The starter is responsible for the milk becoming acid, but the rennet curdles the milk before the bacteria have generated enough acid for clotting.

The milk for making Cheddar cheese is usually pasteurized to kill undesirable organisms and leave the field clear for those in the starter. The pasteurized milk is run into a water-jacketed vat. In a farmhouse the vat may hold only 20 gallons, sufficient to make one small cheese, but in large-scale production some factories use 1,000-gallon vats, one gallon of milk yielding about one pound of cheese. About a gallon of starter is stirred in for every 100 gallons of milk, which is warmed to 29–30° C. (84–86° F.). The milk is allowed to acidify (to roughly pH 6.3) for up to 2 hours before the rennet is added, after which the milk clots in about 10 minutes.

The junket is left undisturbed to become firm for about half an hour, after which it is cut up as evenly as possible into little pieces, each about the size of a pea. Cutting the junket starts its separation into curds and whey; just as when junket served on a dining-room table is cut with a spoon, whey can be seen running from the cut surface. At this stage the vat is almost full of slightly acid whey containing a large number of small pieces of curd which are floating and kept fairly evenly distributed throughout the whey by continual stirring. The growth of the lactic acid producing bacteria continues, and the acidity rises slowly. The temperature of the vat is now raised gradually to bring its contents to about 37° C. (100° F.); this speeds the expulsion of the whey from the curd particles, a process which is accelerated also by the increasing acidity. The pieces of curd tend to sink as they lose moisture, but are kept suspended in the whey by stirring.

After about an hour, stirring is stopped and the particles of curd settle to the bottom of the vat. This is called *pitching*, and many cheese-makers agree that the exact moment of pitching is crucial in making a good cheese. The whey in which the curd is bathed is then run off, checking acid development in the curd as it takes most of the lactose with it, and, moreover, leaves the curd to cool. The whey leaves the vat just as water is let out of a bath, and the curd particles which have matted together at the bottom of the vat are pushed half-way along, away from the whey-exit. The mass of curd is then cut up into blocks and

allowed to drain further. Some photographs which show the making of Cheddar cheese up to this stage are reproduced as Inset 6–8.

At this stage the curd feels rather like rubber, though cheese-makers say that when ready for milling (the next operation) it should be silky in texture and able to be torn into leaves. Milling tears the curd into little pieces, and these are salted, adding about an ounce of salt to every three pounds of curd. The salted curd is placed in metal (formerly wooden) moulds (or 'hoops') which have been lined with cloth (Inset 9). The cheese is pressed under about twenty hundredweight for up to four days, and loses more whey, at least in the early part of the pressing. It is banded and ripened, preferably at about 16°C. (60° F.), for three to six months.

This is only an outline of how Cheddar cheese is made. Numerous variations occur in, for instance, the amount of starter added, the exact length of time for each process, the acidities, temperature, and the way the matted curd is treated. In factories the tendency is to shorten and simplify each process. Compared with Continental varieties of cheese, the most striking characteristics of Cheddar are the high acidity of the curd before pressing and the very considerable pressing to which it is then subjected.

THE BIOLOGY OF CHEDDAR CHEESE

Starter Bacteria

A dairy manager can buy his starters from the dairy department of a university, farm institute, or commercial firm. They reach him either as small bottles of sour milk, or in powder form.

Bacteriologists can use sour milk with a good flavour for making up a new starter; the bacteria which cause the souring may come from poorly washed milking equipment. A minute quantity of the sour milk is spread over the surface of a sterilized bacteriological culture medium (agar) on which colonies of the milk-souring bacteria grow when this is incubated. Parts of individual colonies are picked off into tubes of sterile milk with litmus added; the bacteria sour the milk, turn the litmus pink

and decolorize it by reduction. Preparing fresh individual colonies from the litmus milk culture on more agar increases the chance of getting a pure strain of bacteria. If the bacteria isolated in this way acidify the milk satisfactorily, they are cultured in some more sterile milk, and this sub-culture may be used as a 'mother culture' for cheese-making.

In the dairy, the mother culture is maintained by taking a part of it each day and adding this to a new batch of sterilized milk; this new batch then becomes the mother culture and its predecessor must be discarded. In the course of a day much of the lactose is broken down to lactic acid and, if the culture is left, the acidity eventually becomes so high that the lactic acid bacteria cease to grow. Transferring a few of these bacteria to a new sample of milk gives them a fresh supply of milk sugar and reduces the acidity of their surroundings: it gives them a new chance to survive.

The mother cultures in the dairy are used to make up the bulk starter: a small quantity is added to a churn full of pasteurized milk. After 24 hours this milk is sufficiently acid to be mixed in with the vat full of milk which is to be made into cheese.

Milk souring usually occurs because 'lactic acid bacteria' which are present in the surroundings come in contact with the milk, and, using the milk sugar as a source of energy, reproduce themselves rapidly, spreading through the bulk of the milk. They convert the lactose to lactic acid, which is responsible for the sour flavour and the clotting.

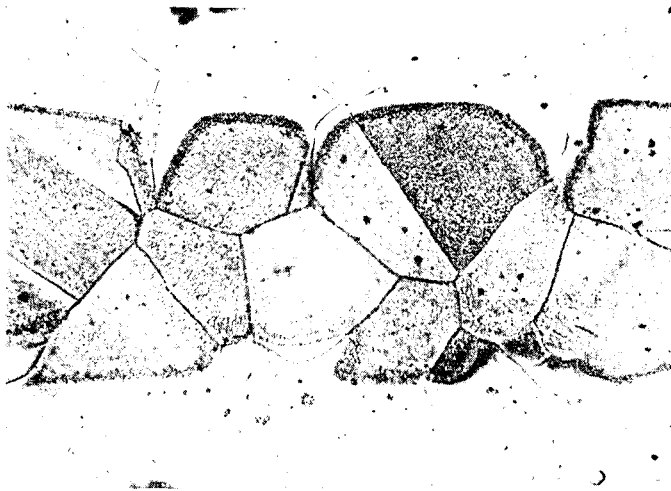
Examination with a microscope shows two kinds of lactic acid bacteria: chains of round (or ovoid) streptococci, called lactic streptococci, and rod-shaped lactobacilli which are also often in chains. The active bacteria in clean-flavoured sour milk are mainly lactic streptococci. Some lactobacilli are usually present, but the acid produced by the lactic streptococci destroys most other types of bacteria in the milk. *Streptococcus lactis* and *Streptococcus cremoris* are the dominant lactic streptococci in ordinary clean sour milk, and although other species are present, they do not produce lactic acid so readily. Most English starters are made up with *Str. cremoris*.

OCEAN-BED PROSPECTING



1. Sono-radio buoy.

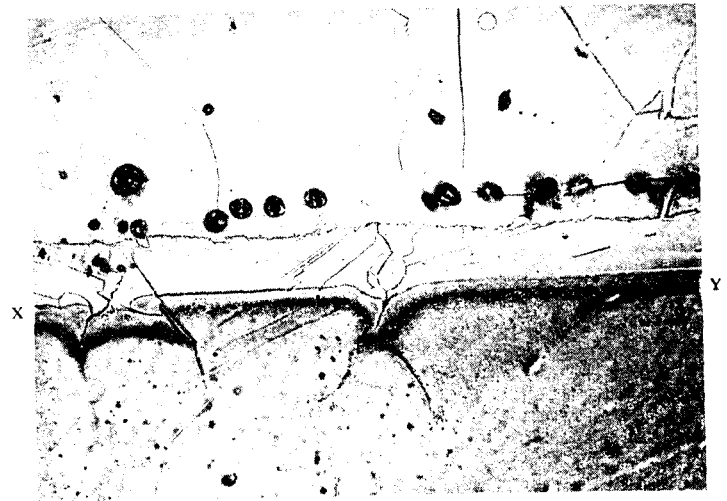
MOVEMENT OF ATOMS IN METALS



2. This and the following photomicrograph show the preferential diffusion of copper along grain boundaries of nickel. Cross-section perpendicular to welded interface. Magnification $\times 290$.



3. As above, but section parallel to the interface, reproduced at the same magnification.

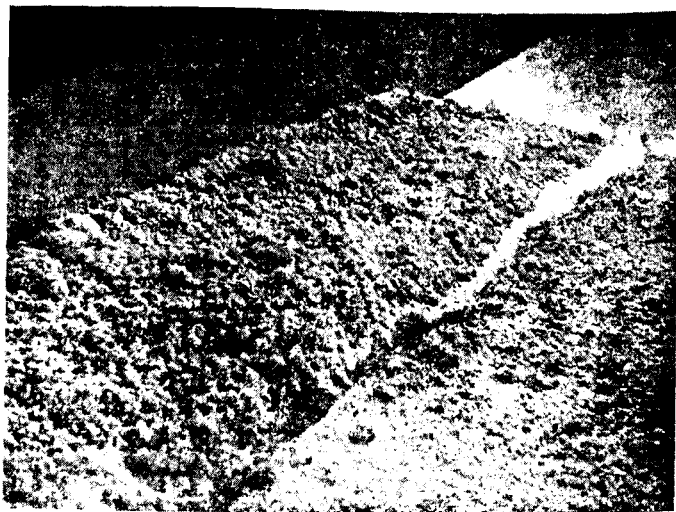


4. This and the following photomicrograph are of copper/nickel couples; they show the extent of diffusion, and the formation of holes on the copper side of the marked interface. Couple annealed at 940°C . Magnification $\times 140$.

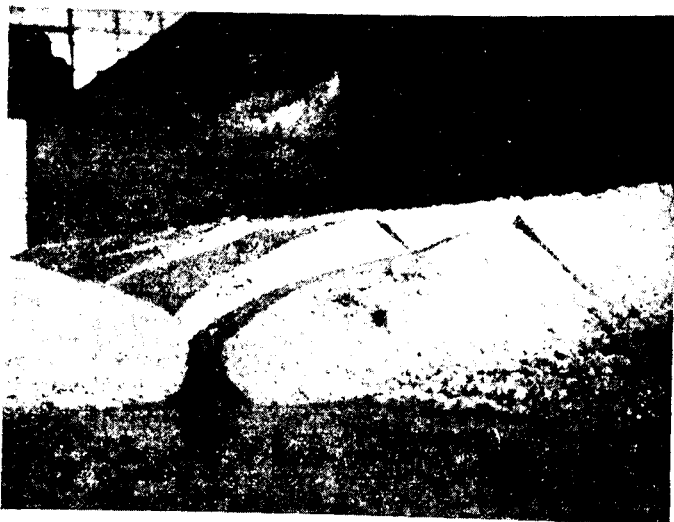


5. As above, but showing couple annealed at 1080°C ., reproduced at the same magnification.

CHEESE BIOLOGY



6. Cheddar: draining the whey off the curd in the bottom of the vat.



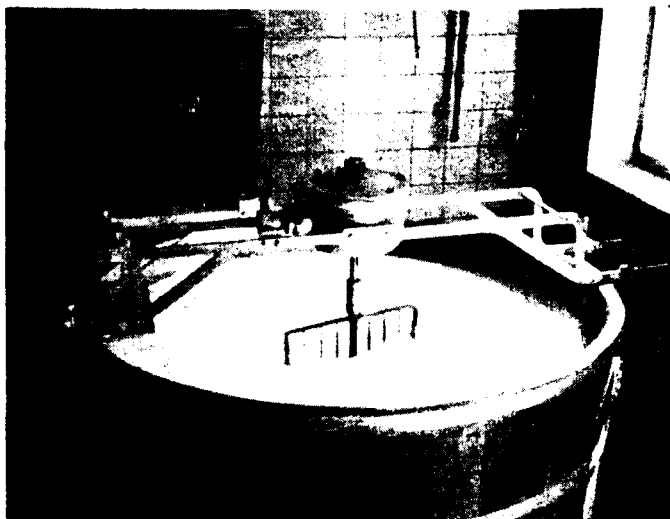
7. Cheddar: the curd in the bottom of the vat is cut into large pieces, and draining continues. (Author's photographs.)



8. Cheddar: the piled curd flattens out under its own weight.



9. Cheddar: after milling and salting, the curd is put into moulds lined with cloths, ready for pressing. (Photograph, N. G. J. Gruber.)



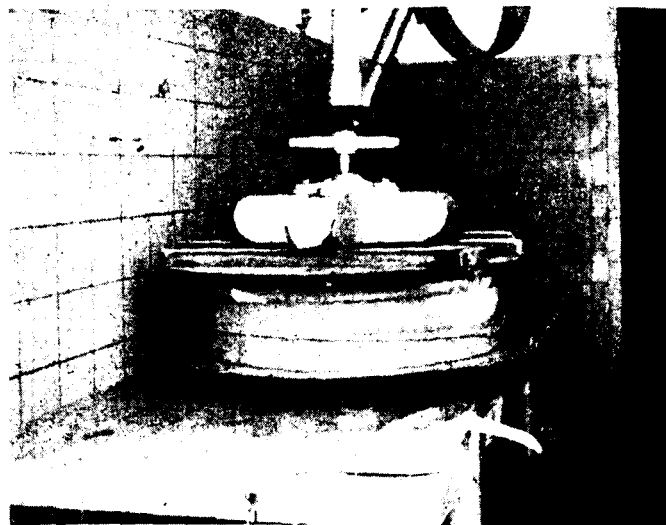
10. Emmentaler: the junket is being stirred in a heated copper vat. (Author's photographs.)



11. Emmentaler: the hot curd is scooped out in a large cloth. Here it is lifted with a pulley, but often this is done by hand.



12. Emmentaler: the cloth full of hot curd is now sitting in a mould (or hoop), in position for pressing.



13. Emmentaler: pressing. Whey is oozing from the curd. (Author's photographs.)

PHASE-CONTRAST MICROSCOPY

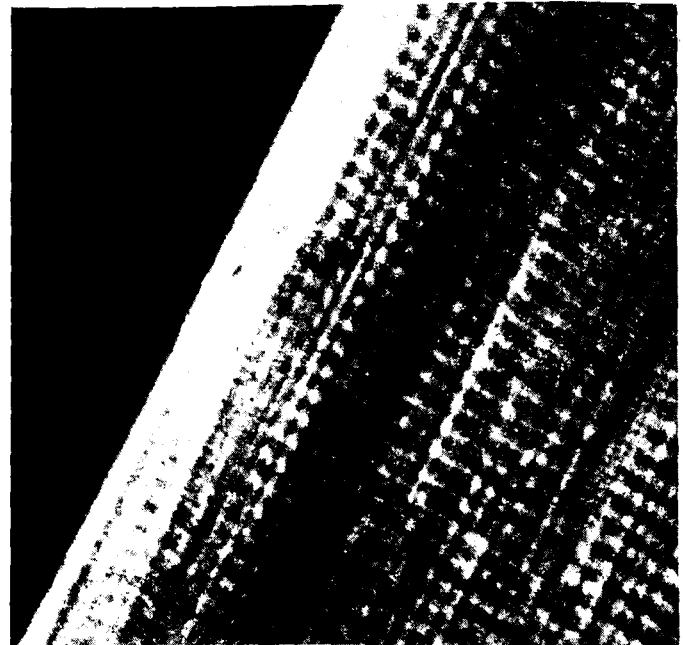
The micrographs on these two pages illustrate the value of phase-contrast microscopy in showing structures in living, intact specimens; they are reproduced by courtesy of Dr R. Barer. (See article by Dr J. G. Thomas in *Science News* 26.)



14. Living, intact salivary gland cell of *Chironomus*, showing nucleus, nucleolus, and giant chromosomes. Magnification $\times 1,200$.



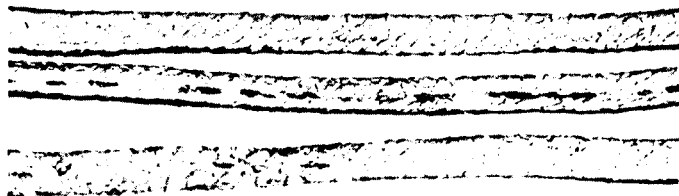
15. Bacteria (*B. megatherium*), showing internal bodies and granules. Magnification $\times 2,200$.



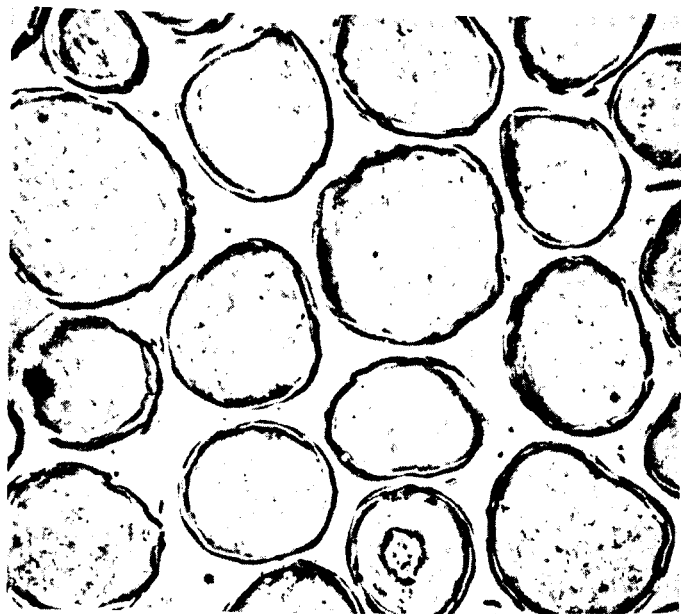
16. Living lizard muscle fibre showing striations. Magnification $\times 1,200$.

TEXTILE FIBRES

The micrographs on this and the following six pages are reproduced from *Identification of Textile Materials* by courtesy of The Textile Institute. There are two micrographs of each fibre: longitudinal, *above*, and in cross-section, *below*.



17. Wool (46s Top), mounted in liquid paraffin. Magnification $\times 140$.



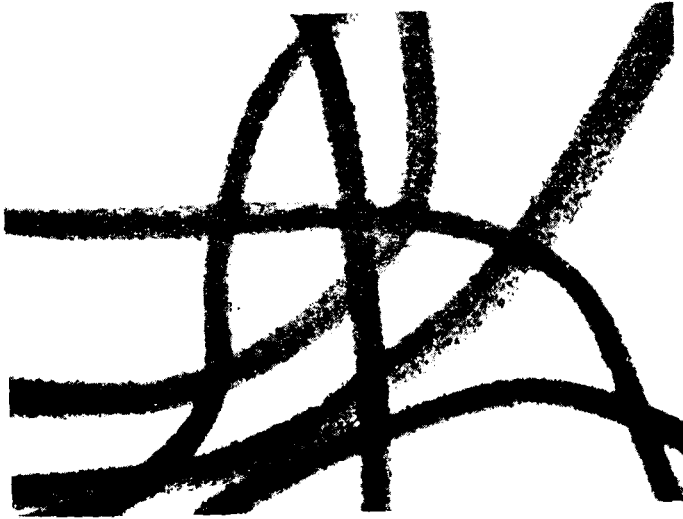
18. Wool (46s Top): in cross-section mounted in 'Permunt'. Magnification $\times 500$.



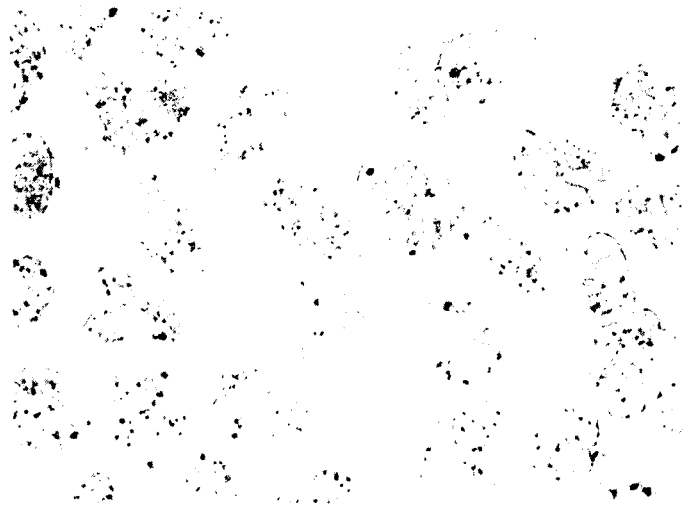
19. Silk, tussah. Magnification $\times 140$.



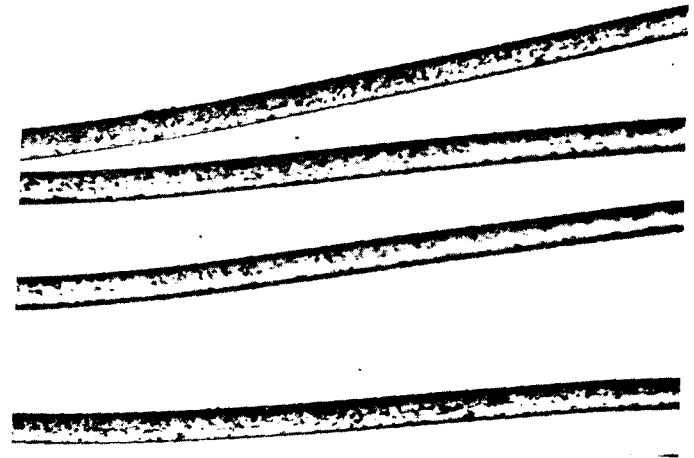
20. Silk, tussah: in cross-section, mounted in gelatine. Magnification $\times 500$.



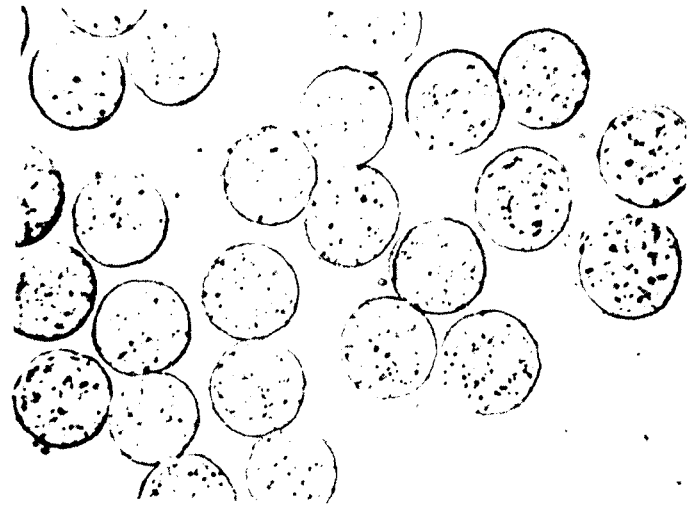
25. Cellulose acetate, phenol delustred, mounted in liquid paraffin. Magnification $\times 180$.



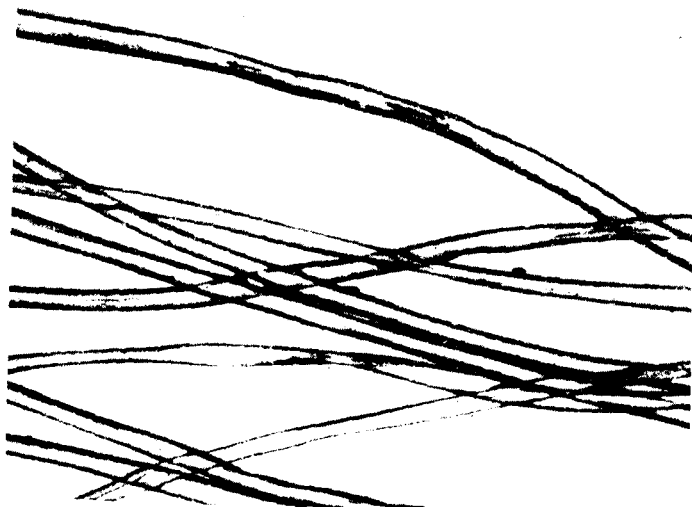
26. Cellulose acetate, phenol delustred; in cross-section, mounted in liquid paraffin. Magnification $\times 500$.



27. Nylon. Magnification $\times 150$.



28. Nylon: in cross-section, mounted in 'Tensol' cement. Magnification $\times 500$.



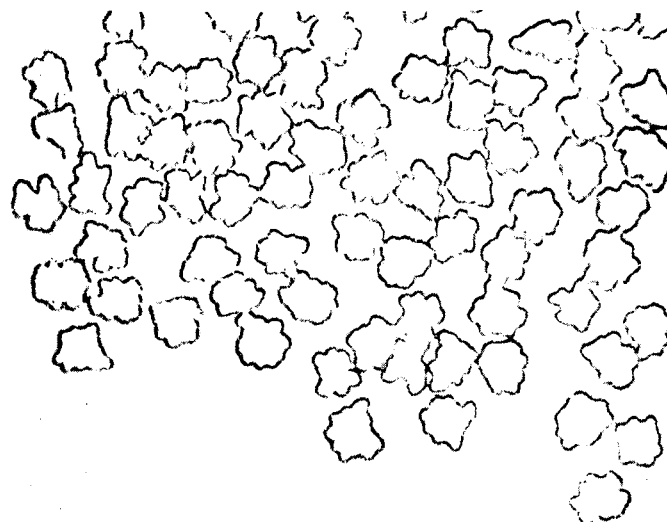
21. Cotton, mercerized, Magnification $\times 180$.



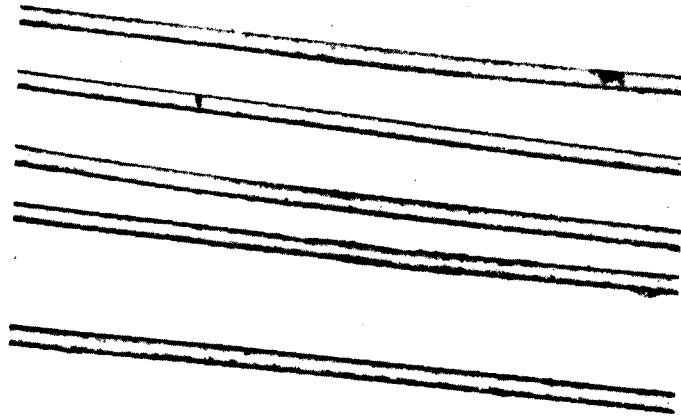
23. Viscose rayon, Magnification $\times 170$.



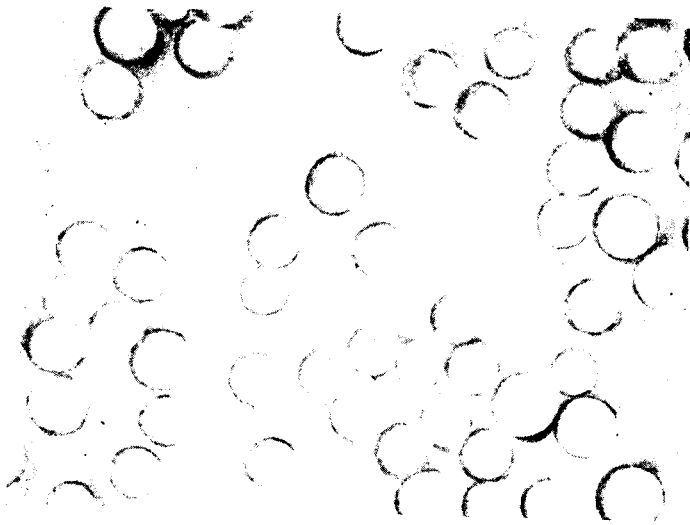
22. Cotton, mercerized: in cross-section, mounted in gelatine, Magnification $\times 500$.



24. Viscose rayon: in cross-section mounted in methacrylate resin, Magnification $\times 500$.



29. Terylene. Magnification $\times 180$.



30. Terylene: in cross-section, mounted in glycerine jelly. Magnification $\times 500$. (The micrographs reproduced on this and the preceding pages are reproduced from Identification of Textile Materials by courtesy of The Textile Institute.)

Cheddar ripening

The starter organisms not only grow and acidify the milk before rennet is added to coagulate it, they continue growth actively throughout manufacture and even when the cheese is in the store. With fermented milk, like kefir milk, the supply of lactose is never used up and the bacteria are killed by the acid they themselves produce, but in Cheddar cheese most of the lactose is run off in the whey, and the buffering capacity of the curd is sufficient to compensate for most of the acid which can be produced from the limited amount of lactose remaining in the cheese.

The acid speeds up the action of the rennet in clotting the milk; it helps, also, in expelling whey from the curd particles which contract more under acid conditions and, in addition, mat together better. Another important role of the acid is to protect the cheese against the growth of undesirable organisms which spoil the texture, flavour, or appearance, or make it dangerous to eat. These unwanted bacteria cannot flourish in acid conditions and in the presence of salt, and so a well-made Cheddar may remain in excellent condition for many months if it is kept cool. Moulds often grow on the surface of the rind, but if the cheese is properly made, they cannot penetrate the firm curd where there is little oxygen. Lastly, acid is necessary for the action of pepsin. This protein-splitting enzyme from the stomach is included in rennet and attacks the casein.

Lactobacilli, as well as lactic streptococci, seem to be important in the ripening of Cheddar cheese. Probably in cheese, as in milk, *Str. cremoris* can grow more rapidly than the lactobacilli, so their growth is limited to when the lactose is used up and conditions are unsatisfactory for the streptococci. Without lactose, the lactobacilli probably use the protein derivatives which result from the action of the pepsin on casein, and so carry the protein breakdown a stage further. The change of the tough, rubbery, unripened curd to the consistency of mature cheese is symptomatic of protein breakdown.

Later, the lactobacilli, too, die off; and a recent examination

of two 7-year-old Cheddars showed that there were indeed very few bacteria left.

The lactobacilli are thought to be responsible for breaking down some of the fat by means of an enzyme called lipase, liberating fatty acids. These fatty acids, freed from the fats of which they form a part, are likely to have considerable influence on the flavour of the cheese. Lipase produced by bacteria has been detected after about a week's ripening: some lipase is present in the original milk before bacteria invade it, but this lipase, unlike that of bacterial origin, is not active at the acidity found in cheese even if it survives pasteurization. Now the lactobacilli isolated from cheese do not appear to attack fat during their normal life cycle, but old cultures of these lactobacilli, where there are large numbers of dead (autolysed) cells have been found to break down milk fat. Lipase has in fact been released from cells of a strain of *Lactobacillus casei* by digesting away their cell walls. It is known that butyric acid does not appear in Cheddar until the cheese is mature: this later appearance may be due to liberation of lipase from disrupted lactobacilli when conditions become unsuitable for their further growth.

Two workers at Wisconsin report that various mixtures of amino-acids in the proportion they found in cheese have distinctive cheese flavours. Another two Americans have tried to study the changes in flavour which result from adding specific amino-acids to the curd after salting. They concluded that perhaps serine gave a marked improvement; this may not, however, be a direct effect on flavour. More likely, serine favours growth of the lactobacilli which produce the flavour-giving substances. Quite independently, two British bacteriologists have suggested that serine is 'an amino-acid of particular significance in bacterial metabolism'.

Experiments at Cornell using a strain of *Streptococcus faecalis*, isolated from human faeces, in starter have led to the suggestion that this organism is important for producing a good Cheddar flavour. The assumption is that *Str. faecalis* is normally present in raw milk and is the main organism responsible for

producing a substance called tyramine, derived from the amino-acid tyrosine. The Cornell workers found a close relation between tyramine content and flavour, though tyramine added to the cheese did not give the Cheddar flavour. They concluded that the amount of tyramine indicated the extent of activity of *Str. faecalis* and probably of other bacteria, such as lactobacilli. It will be interesting to know if, in this case, English palates correspond with American taste, and also whether or not *Str. faecalis* in cheese is capable of producing ill health.

The characteristic flavour of mature Cheddar cheese must be produced by the combined effects of the breakdown products of casein, milk fat, and lactose. These products include amino-acids and ammonia, volatile fatty acids, esters, aldehydes, and ketones. Figure 1 is an outline of the ways in which they may be produced:

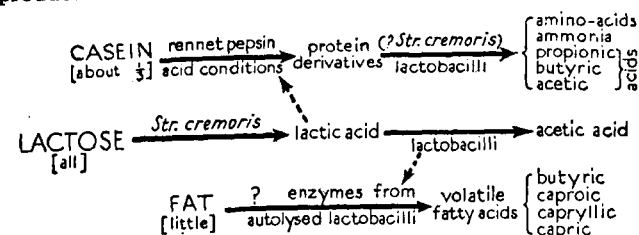


Fig. 1. The breakdown of casein, milk fat, and lactose in ripening Cheddar cheese.

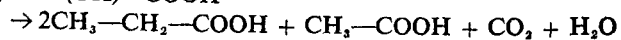
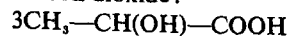
THE BIOLOGY OF OTHER CHEESES

Emmentaler and Gruyère

Emmentaler (Swiss) and Gruyère (chiefly French) are very similar, and perhaps the most important Continental hard cheeses. They differ greatly from Cheddar, particularly in never becoming very acid. The curd is heated in a deep round vat to over 50° C. and put in the press when still hot. These are large cheeses which can maintain, therefore, for a long time, the high temperature which they have reached in the vat, even when they are in the press. So they are distinguished from most other

cheeses by having starter organisms which thrive at relatively high temperatures; in technical language, a thermophilic lactic flora. This is a good example of how size has become characteristic for a cheese: a small Emmentaler would cool rapidly and the starter organisms would not flourish. Some sort of cheese would result, but it would not be an Emmentaler. Typically, a Cheddar may weigh 50 to 60 lb., but a Gruyère must be between 70 and 120 lb., and Emmentalers weigh anything from 160 to 230 lb. For mechanical and other reasons, these huge cheeses cannot be tall and cylindrical like a Cheddar, but are flat and disk-like. J. B. S. Haldane has pointed out that a giant man sixty feet high would have bones too small in area of cross-section to carry his increase in weight, and that few insects are more than half an inch thick as their tracheal system is not capable of giving an adequate oxygen supply to a large animal. 'For every type of animal', he says, 'there is a most convenient size, and a large change in size inevitably carries with it a change of form.*' This is true of cheeses too.

Swiss cheese starters may contain four different types of bacteria which are important for normal ripening. *Str. lactis* (or *Str. cremoris*) starts the acidification, as in Cheddar-making, but when the temperature rises above 37° C. *Streptococcus thermophilus* takes over. This is the main acid-producing organism, and it flourishes between 40 and 45° C. Lactobacilli are thought to do much the same for Swiss cheese as for Cheddar, continuing protein breakdown begun by the rennet and being responsible for the greater part of flavour development. The fourth kind of organism, the propionic acid bacteria, converts lactic acid produced by the streptococci into propionic acid, acetic acid, and carbon dioxide:



This carbon dioxide is responsible for forming the characteristic 'eyes' or holes. Well-formed spherical eyes depend on the elasticity of the cheese: thus eye-size is not so much dependent on the number of living propionic acid bacteria in the cheese as it

*Possible Worlds. London, Chatto & Windus (1927).

is affected by acidity, salt concentration, and, in particular, temperature. For normal eye development, an Emmentaler is taken from the cool store when two weeks old and kept for a few weeks at 20–22° C. The corresponding temperature for Gruyère, on the other hand, is 16–17° C., and it is for this reason that Gruyère eyes are smaller than those in Emmentaler. A Gruyère cheese kept at an Emmentaler temperature develops Emmentaler-size holes. But the eyes are larger in Swiss than in other comparable cheeses such as Dutch Gouda, mainly because the high temperature mats the curd together to such an extent that diffusion of gases within it is considerably reduced, and bubbles can be formed. Inset 10–13 show the making of Emmentaler, and Figure 2 summarizes some aspects of Swiss cheese ripening.

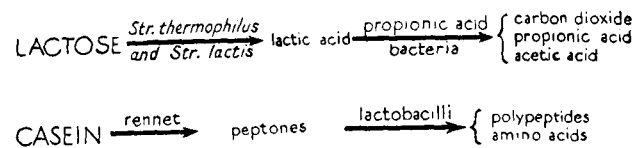


Fig. 2. The ripening of Swiss-type cheese.

Stilton

The main ripening agent in Stilton cheese is the fungus *Penicillium roqueforti*. This species forms the characteristic 'blue veins', not only of Stilton, but also of French Roquefort and Italian Gorgonzola, although different methods are used in the manufacture of each of these types of cheese and they differ markedly both in flavour and in appearance.

P. roqueforti does not develop in Stilton before the cheese is about two months old, and until this stage is reached ripening is akin in essentials to that of Cheddar. E. R. Hiscox and her colleagues at the National Institute for Research in Dairying (N.I.R.D.), of Reading University, describe three stages of ripening for Stilton: the *white stage* before the mould begins to grow; the *period of active mould growth*; and *maturation*, a stable period when the mould has stopped growing fast.

The chief characteristic of the white stage is an increasing acidity. Although, in contrast with Cheddar, little or no starter is used and a Stilton begins life as a low-acid cheese, the curd stands in its whey for several hours, a condition favourable to acid production as the lactose remains available for fermentation. Moreover, the curd is not heated and is wet when placed in the hoops; so, as there is no pressing, drainage is very slow. The acidity increases gradually over several days, so that a young Stilton is more acid than a Cheddar. E. C. V. Mattick of the N.I.R.D. has shown how calcium and phosphorus pass into the whey, the calcium content of which steadily rises as the acidity increases. In a ripe Stilton much of the remaining calcium is present in the cheese coat (which most people throw away) probably because evaporation of whey occurs from the surface of the cheese leaving behind the salts which had been dissolved in it. Thus, unlike most other varieties of cheese, Stilton is a poor source of calcium and phosphorus. Whereas Cheddar contains 58 per cent of the calcium present in the original milk, and 50 per cent of the original phosphorus, the corresponding figures for Stilton are 7 per cent and 27 per cent respectively.

After a month or two, the blue mould begins to grow, and active growth continues for about six weeks. During this period, the *Penicillium* is responsible for a far greater fat breakdown than occurs in Cheddar or Emmentaler, and it attacks also the protein. As a result of the protein breakdown there is a great increase in the amino-acid content of the cheese; and the amino-acids in their turn are attacked, liberating ammonia. Compared with Cheddar, in Stilton normally about twice the amount of nitrogen is in the form of ammonia. The considerable fat-splitting activity produces a cheese with relatively large amounts of volatile acids, and these can be separated by steam distillation. At the same time, carbon dioxide accumulates from these metabolic activities, just as living human beings give off carbon dioxide, and the cheese becomes less acid. Nowadays growth is hastened by piercing the cheeses with skewers or needles: this reduces the concentration of carbon dioxide as it can escape more easily through the holes, and more oxygen is available to

the mould from outside. Heavy pricking lets in far too much oxygen, and this results in a rapid growth, so that there is too much breakdown of fat and protein, and the subtle flavour is not produced.

The period of maturation begins when the mould's active growth is finished. Hiscox and her colleagues describe this as a more stable period when volatile acids tend to accumulate and there is a mellowing of texture and development of flavour.

Roquefort and Gorgonzola

Authentic Roquefort cheese is made from sheep's milk, but its course of ripening is fairly similar to that of Stilton. Much of the work on blue-veined cheeses has been carried out with imitations of Roquefort, both in the United States and in Denmark, cow's milk being used instead of ewe's.

In trying to discover why *P. roqueforti* is dominant in blue cheese, and not other *Penicillia* or *Aspergillus*, it was found that *P. roqueforti* tolerates a high concentration of carbon dioxide, a low oxygen tension and a high concentration of salt better than any of the others. In spite of the hardness of *P. roqueforti*, these blue-veined cheeses are limited in size, as it is difficult for the mould to penetrate a really large cheese sufficiently rapidly. Gorgonzola, the largest, does not usually weigh more than 20 lb., while a Roquefort may be as small as 5 lb. Traditionally, Stilton and Gorgonzola receive their *Penicillium* by chance infection, whereas Roquefort is deliberately inoculated. This means that the two former cheeses are the slower to ripen.

A number of different suggestions have been made as to what gives a typical Roquefort flavour. Probably some ketones are important, and these seem to be produced from the volatile acids (which are liberated from the fat) by the action of the *Penicillium*. So *Penicillium* seems not only to break down the fat and liberate volatile acids, but to take the breakdown a stage further, attacking some of the volatile acids themselves. The mould is active for longer in Roquefort than either of the other two kinds of cheese by virtue of its initial inoculation; it has a smaller volume of cheese to attack and so ripening goes a good deal

further, and butyric acid is present in much larger quantities.

A summary of the ripening of blue cheeses is given in Figure 3.

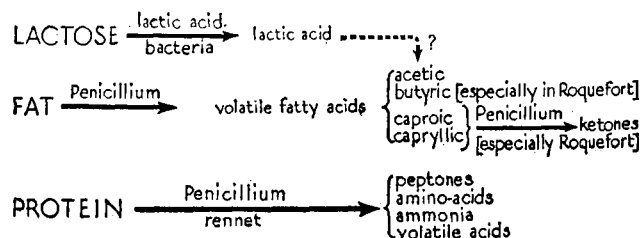


Fig. 3. The ripening of blue cheese.

Camembert and Brie

Camembert and Brie are soft cheeses, and they must be small because they are ripened from the outside by organisms which grow on the surface. They are rennet cheeses, acidified by lactic acid bacteria, but some yeast settles on the surface and probably lives on the casein and lactic acid. Anyway, it dries the surface and prepares the way for *Penicillium camemberti*, a white mould which appears on the surface two or three days after making. The *Penicillium* forms a surface felt, enzymes from which diffuse into the cheese so that ripening starts from outside. The mould reduces the acidity and attacks the protein. Two or three weeks later the *Penicillium* is succeeded by a surface bacterial growth, usually the red-coloured *Bacterium linens*, and this is responsible for the greater part of the protein breakdown.

Gammelost

In Norway Gammelost has much the same status as has Stilton in this country. This cheese ripens both from the outside like Camembert and also from within in the manner of the blue cheeses. It is made exclusively from skim-milk, and so the ripening is almost entirely protein breakdown. *Mucor* is the mould which grows on the cheese coat, and *Penicillium roqueforti* is introduced inside. Like other Norwegian cheeses Gammelost is very strange indeed. It is dry, flaky and inclined to be bitter,

probably from the protein derivatives. Although the finished cheeses are very similar, there are two main methods of making which have arisen in two different localities, Sogn and Hardanger Fjords. In both cases starter is added and the milk left for about two days to become very sour; coagulation is by acid alone, and no rennet is used. By the Hardanger method the soured milk is heated to the very high temperature of 60° C., but by the Sogn method it is boiled. The curds are dried and infected with the correct fungi. If fat is left in the milk from which Gammelost is made it gives unpleasant flavours because the moulds attack the fat instead of concentrating on the proteins. On the other hand, Mysost, the most popular Norwegian cheese, can be made from whey and cream, with no curds!

CHEESE IN GENERAL: A COMPARATIVE SUMMARY

Type of milk

Cheese-making provides a pragmatic division of milk into two classes; cheese can be made from a milk which forms hard curds when treated with rennet; it cannot be made from a milk which forms no curd, or only a very soft one. In Europe, cheese is made with milk from cows, ewes, goats, and reindeer. But, although mare's milk is consumed in parts of Eastern Europe, for instance

TABLE 1: Approximate Chemical Compositions of some types of Milk (as percentages to nearest whole numbers)

Type of Milk		Water	Fat	Protein	Sugar
Cheese-making Milks	Cow	87	4	3	5
	Ewe	80	7	5	5
	Goat	82	5	3	5
	Reindeer	68	17	10	3
Other Milks	Mare	90	2	2	7
	Ass	90	1	2	7
	Human	87	4	2	7

(Sources: Davies, W. L., *Chemistry of Milk*, 2nd edn, London, Chapman & Hall, 1939; and Davis, J. G., and Macdonald, F. J., *Richmond's Dairy Chemistry*, 5th edn, London, Griffin, 1953.)

as the fermented milk drink kumiss, it cannot be used for cheese; in common with asses' and human milk it does not form a hard curd. These soft-curd milks contain more water and sugar than the cheese-making milks, and they are lower in protein.

This distinction is important in quite another practical way. It is generally thought that human infants should not be given plain cow's milk undiluted and untreated, lest it should form hard clots in the infant's stomach causing discomfort, a great deal of the curd appearing undigested in the faeces. Dilution with water and boiling both go to soften the curd. This can be seen very easily in the kitchen if boiled milk is used in an attempt to make junket: no proper curd is formed.

Treatment in the vat

A sour milk curd is tough and crumbly in contrast with rennet curd which is soft, markedly elastic, and able to shrink, thereby eliminating whey, just as serum separates from a retracting clot of blood.

Temperature is crucial in treating the vat of curds-and-whey, and it is important in characterizing the cheese. Stilton, for instance, undergoes no heating in manufacture and so draining the curds is slow and a high acidity develops gradually over several days. On the other hand, heating to 37° C. when making Cheddar ensures a relatively rapid expulsion of whey; while the 50° C. to which Emmentaler curd is raised gives the curd its peculiar plasticity essential for the formation of proper eyes, and in addition favours the growth of a different type of micro-organism.

Ripening

Although temperature partly determines what organisms can grow in the cheese, nowadays the bacteria and fungi are usually introduced artificially. The lactic acid bacteria seem to be important for nearly all kinds of cheese. In lactic cheese, Cheddar and Emmentaler, they are the *only* normal ripening organisms (if Emmentaler's propionic acid bacteria are included in this category), while in the blue cheeses and Camembert they are only the first of a series which break down the milk constituents: in

all ordinary cheeses lactic acid bacteria are the *first* active organisms.

The chief feature of ripening is the conversion of insoluble proteins into soluble substances. In the Swiss and Dutch types of cheese usually only about one-third of the casein is converted into water-soluble compounds, whereas in soft cheeses like Camembert nearly all the casein is broken down. Rennet is responsible for most of the initial attack on the protein, and the micro-organisms carry the breakdown further. As the action of cheese-ripening bacteria tends to be inhibited, while that of rennet is promoted by acid, the Swiss and Dutch cheeses which are hard but not very acid contain more amino-acids than the soft cheeses: the breakdown is less extensive, but more thorough.

Fat is broken down more rapidly in a cheese, like Roquefort, permeated with moulds. In contrast, fat hydrolysis is very slow in the hard, rich cheeses and they require a long ripening period to develop a mature flavour.

Lactose, on the other hand, is the first major constituent to disappear in cheese ripening. Normally all the lactose is broken down to lactic acid which may be metabolized further, as in Emmentaler where it is converted into propionic acid.

Both the blue-veined cheeses, Stilton, Roquefort, and Gorgonzola, and also Cheddar and the Swiss and Dutch varieties of hard cheese, are ripened from within. For Camembert and Brie ripening starts on the surface and works inwards. And not only is this external attack able to ripen the whole cheese because of its small size, but also because the flat shape of the cheese, too, provides a greater surface area (relative to the volume) from which the attack may be made.

SOME MALADIES OF CHEESE

French dairy scientists distinguish between *les défauts de fabrication* and *les maladies du fromage*. Here their example is followed, and faults in manufacture are considered only where they cause a cheese malady of microbiological origin. Such maladies illustrate further the individuality of the different kinds of cheese. Not only do cheeses have their characteristic bacteria

in normal development, they are characterized also by the types of infection to which they are liable.

Trouble from alien bacteria in Cheddar cheese is usually symptomatic of insufficient development of acid by the starter streptococci. In these conditions certain coliform organisms, such as *Bacterium aerogenes* or sometimes *Bacterium coli*, can grow rapidly. *Bact. coli* is the most common intestinal bacterium, and *Bact. aerogenes* is very similar, but is more usually associated with soil and plant materials. Neither is normally responsible for human disease, but for a cheese they can be very serious indeed. They flourish without oxygen, and break down lactose or lactic acid generating carbon dioxide. This gas produces a large number of tiny holes in the curd, and causes it to float in the vat instead of pitching. The curd is then difficult to handle, and the holes spoil the texture of the cheese.

The same organisms cause trouble in Swiss cheese too. Here the holes are formed when the cheese is in the press or some days after. They can be detected in the finished cheese by tapping it with a finger, just as a doctor taps a patient. Eventually the gas produced may bulge the young cheese, and often the body will be leathery and have a sharp disagreeable flavour. The effect on flavour is most marked if lactic acid, as opposed to lactose, has been fermented by the coliform organisms.

Trouble with *Bact. coli* and *Bact. aerogenes* is best prevented by cleanliness. An active starter should hold back their development, not only by high acid-production, as in Cheddar-making, but also because of antagonism between the lactic acid bacteria and the coli-aerogenes group. But in factory cheese-making the emphasis on speed may lead to insufficient time being given for the starter bacteria to produce enough acid. Finally, pasteurization, now usual in making Cheddar cheese, is effective against coliform organisms, although it decimates the lactobacilli of the milk too, and may very well contribute to the poor flavour of present-day Cheddar. It certainly has very serious consequences for the Swiss type of cheese in holding back the growth of the propionic acid bacteria.

There is a contrast between Cheddar and Swiss-type cheeses

in liability to infection with another type of organism, the butyric acid producing bacteria of the genus *Clostridium*. These organisms not only flourish without oxygen, but also form resistant spores which survive pasteurization. They are held in check by salt, particularly if the acidity is high, as in Cheddar cheese. In Emmentaler and Gruyère, on the other hand, not only is the acidity less, but the salt is merely rubbed into the rind, instead of being fairly evenly distributed through the cheese. Conditions are thus more favourable to these organisms, and they cause serious losses in the Swiss cheese industry. They break down lactic acid into butyric acid, hydrogen, and carbon dioxide. The butyric acid gives the cheese an unpleasant flavour: the hydrogen produces large holes, blowing the cheese out; while the carbon dioxide tends to go into solution and does not give much trouble. These big holes are often separated by only a thin membrane; and, if one of these 'blown' cheeses is pierced with a skewer, and the gas squeezed out by pressing with the hand, the escaping hydrogen can be ignited, burning with a steady flame.

Cheese Starter Troubles

Failure of the lactic streptococci, used in starters, to produce enough acid for Cheddar cheese making may be caused by bacteriophage (or phage), a virus parasitic on the bacteria in which it multiplies. It can bring about a rapid destruction of the bacterial cells, so that no more acid is produced. The first report of a starter failure from this cause was given in 1935 by H. R. Whitehead and G. A. Cox of the New Zealand Dairy Research Institute. A pure culture of *Str. cremoris* in commercial use as a cheese starter sometimes failed to develop acid. Whitehead and Cox had found previously that aeration of a starter sometimes led to a sudden failure. Although the mother culture of this batch of starter appeared normal, vigorous stirring of sub-cultures induced loss of activity.

They set out to study the differences between the active and inactive batches of starter culture. They added a few drops of each to tubes of warm pasteurized milk which had been coloured with methylene blue. The blue colour in both tubes disappeared

in two or three hours, and this suggested that the streptococci in each were satisfactorily active. The methylene blue test on a milk sample measures one type of bacterial activity in it; the colour generally disappears more rapidly the higher the numbers of bacteria. Microscopic examination confirmed that there was a normal growth of streptococci in both samples. But an hour later the blue colour had returned to the sample of aerated starter, and no bacteria could be seen with the microscope. Growth went on normally in the active starter.

Whitehead and Cox concluded that some agent had destroyed the streptococci in the inactive culture and that this destruction accounted for the failure in acid production in the cheese vat. They found, moreover, that minute quantities of milk from the tube in which the streptococci were destroyed were sufficient to cause similar destruction in an active culture.

Although phages had been observed for most types of bacteria, they had only once before been noticed attacking lactic streptococci. A final observation made Whitehead and Cox certain that they had found a phage attack on the starter organisms. A culture which had become inactive after aeration was grown on agar. Normally the bacteria grow fairly evenly on the agar, but in this case there were large numbers of clear circular patches free from growth. These patches are characteristic of bacterial cultures attacked by phage.

Whitehead and Cox showed later that droplets of moisture carry phage, and it is believed to remain dormant in dust particles for a long time. Factories suffering from chronic phage-attack have been saved from stoppage by making up bulk starter in a special room away from the main part of the factory, and by taking thorough precautions against the introduction of infection into this room. The best way to avoid trouble is to use starters made up with strains of lactic streptococci which are attacked by different phages, the latter being highly, though not completely, specific.

Less disastrous starter failure may be caused by inhibitory substances produced by certain types of bacteria growing in the milk. These substances inhibit the growth of other micro-organ-

isms: that is, antibiotics may be produced by bacteria in the milk. It was Whitehead, again, who in 1933 found that a strain of *Str. lactis* with this characteristic was responsible for a starter failure. Whitehead's inhibitory strain was passed on to England, where E. B. Anderson and L. J. Meanwell of United Dairies isolated strains of *Str. lactis* resistant to the inhibitory substances. These strains have been used for commercial cheese-making over a number of years without losing this resistance. Just after the war, A. T. R. Mattick and A. Hirsch of the N.I.R.D. described a strong antibiotic, nisin, produced by descendants of Whitehead's strain. Nisin has been found to be effective against tubercle bacilli and organisms responsible for mastitis. It has been used by Hirsch and his colleagues to tackle the problem of 'blowing' in Swiss-type cheese caused by clostridia. A nisin-producing strain of *Str. lactis* introduced into the starter was found to prevent entirely the destructive activities of the clostridia without impairing the normal starter activity.

Lactobacilli which produce inhibitory substances have been isolated also from Gruyère cheese. Hirsch and his co-workers set out to find whether naturally produced antibiotics within the cheese could account for the absence of butyric acid fermentation despite the presence of large numbers of clostridia. They found that some of the lactobacilli in young Gruyère cheese do produce large quantities of inhibitory substances which have different characteristics from those of nisin but resemble hydrogen peroxide; they are probably important in suppressing the butyric acid fermentation.

Penicillin itself, now used for the treatment of mastitis in cattle, has caused trouble in cheese making. Being soluble in water, it gets from the udder into the milk, and later may inhibit the growth of starter bacteria. Moreover, pasteurization does not destroy penicillin. The quantity needed to produce an inhibitory effect is too small to be detected by any readily applied test; and a small dose of penicillin is not easily distinguished in practice from a light infection of the milk with bacteriophage. Penicillinase, an enzyme which destroys penicillin, has been used experimentally to get rid of penicillin in milk; but

its use in cheese making may be open to a number of objections. It does, however, now seem to be practicable in some countries, and is the only really effective remedy. The attempt to restore Cheddar flavour to American cheese by using *Str. faecalis* is of interest in this connexion since *Str. faecalis* is resistant to penicillin.

CONCLUSION

Cheese-making is an ingenious technique for persuading harmless micro-organisms to work in an organized way towards preserving most of the nutritional value of milk in a palatable form which can be transported and stored more easily than liquid milk, and in which it is far more difficult for disease-producing bacteria to multiply. Weight for weight, cheese compares very favourably as a food with grilled beef-steak, cheese having less water and as much protein. It is a better source of calcium and phosphorus than beef, though it contains less iron.

The scientific study of cheese-making has advanced most in Denmark and the U.S.A., where European cheeses are imitated with varying degrees of success, and recently in New Zealand, where the best factory Cheddar is made. Research has influenced the industry greatly. Before 1900, for example, starter cultures were not used at all for cheese-making. Now, as a direct result of the work of Storch in Denmark and Conn in the U.S.A., starters are almost universal and their use goes a long way towards making very large numbers of cheeses marketable. Bacterial infection cannot be left to chance when cheese is produced on a large scale to meet the needs of a great urban population.

Even though those who have achieved the greatest advances in cheese science have been academically trained, they have been concerned, in Orla-Jensen's words, with 'constantly clearing up more and more of the problems which confront the practical dairyman in his daily work'. But not only have scientists affected dairying, dairy research has furthered academic science. Orla-Jensen's work with the lactic acid bacteria helped to reorientate bacteriologists from morphology towards considering *activity* to be the most important aspect of bacteria. Cheeses are good

material for the study of microbiological population changes, and starter difficulties have stimulated research on aspects of bacteriophage and antibiotics of an importance far beyond dairying.

Considerable differences from one individual cheese to another, and even between two different parts of the same cheese, present research difficulties which may have to be overcome by using large numbers of cheeses. As the important thing about cheese is whether it is good to eat, another difficulty comes from the necessity for assessing quality. With many hundreds of cheeses grading is very difficult, yet Ministry of Food officials attempt it daily. Not long ago most cheese-makers estimated milk temperature with their elbows and acidity by smell. Now they use a thermometer, and a somewhat unsatisfactory titration for 'per cent lactic acid'. Scott Blair of the N.I.R.D. has been studying the psychology of subjective judgements in cheese-making and cheese judging, with a view to expressing them in terms of physical measurements.*

Some of the finest cheeses, Double Gloucester and Port Salut for example, have been made for a small clientele. But nowadays cheeses must be produced for large numbers of townfolk and in these circumstances must be necessarily exoteric in flavour, not for the gourmet. For a gastronome there may be one day, and one day only, on which a cheese is right for eating; but for most people a piece of cheese must last at least several days.

A widespread demand for good cheese could be helped by scientific work. A detailed knowledge of chemical and biological changes in ripening makes factory control possible, and factory methods can be changed to give better results. New kinds of cheese can be developed to suit factory conditions: Danish-Swiss (Samsö) cheese is an example of such an invention. The Danes, with no traditional cheese of their own, have imitated cheeses from France, Italy, Switzerland, Germany, and Norway, but not England. It is probably relevant that English cheeses are

*He described some of this work in 'Measuring Craftsmanship', *Science News* 7 (1948).

generally very acid whilst, on the whole, Continental ones are not. The low acid cheeses are in many ways better suited for factory production; for one thing labour time is not wasted awaiting acid development. Perhaps a suitable low acid cheese could be developed for British manufacture, more economical with labour than most traditional English cheeses.

But even today not nearly enough is known about cheese biology; there is, for instance, much to learn about the microbiology and chemistry of ripening. Many of the problems to be solved are linked with academic questions of, for example, the distinction between two types of organism, such as *Str. cremoris* and *Str. lactis*. These bacteria are so much alike that many authorities do not distinguish between them. In the scientific study of cheese, it is important to know exactly what organisms are present and how they differ from those in other cheeses; establishing the identity of organisms is still a major problem in bacteriology. The interesting differences are those of their chemical activity: whether they break down lactose to lactic acid, attack proteins, or produce toxic substances. If there are many intermediates between the different forms, and if one population of bacteria can adopt the characteristics of another, they are going to prove very difficult organisms to sort out. This is just one example of many biological problems arising in connexion with cheese research.

Those who remember good pre-war English Farmhouse Cheddar ask what right has present-day factory cheese to the same name. They ask, too, if a really good factory Cheddar can in fact be made. The Danes have succeeded better than anyone in making factory cheeses with good flavours. But Denmark has not undergone an intensive industrial revolution, and perhaps in consequence Danish taste has suffered less than English. No amount of dairy science can create good factory cheese for a people without gastronomic education and who have for many years been unable to afford really good quality food. As a consequence they do not demand high standards, and factory owners emphasize increased output and high yields, rather than producing the finest quality cheese.

(I wish to thank the following for much helpful criticism; Dr A. Hirsch, Miss E. R. Hiscox, Professor H. D. Kay, F.R.S., and Dr G. W. Scott Blair.)

FURTHER READING

The following are general references. A few more specialized publications are listed below under subjects:

Farm and Creamery Cheese-Making. 3rd edn (amended reprint). London, H.M.S.O. (1947).

Orla-Jensen. *Dairy Bacteriology*. 2nd edn. London, Churchill, 1931.

Hammer, B. W. *Dairy Bacteriology*. 3rd edn. New York, Wiley, 1948.

Dorner, W., Demont, P., and Chavannes, D. *Microbiologie Laitière*. 2nd edn. Lausanne, Librairie Payot, 1945.

On English cheeses (*Journal of Dairy Research* in each case):

Dearden, D. V., and others (nutrients in English cheeses), **14**, 100 (1945); Hiscox, E. R., Harrison, J., and Wolf, J. Z. (Stilton), **18**, 296 (1951); and Mattick, E. C. V. (also Stilton), **18**, 305 (1951).

On Swiss Emmentaler and Norwegian Gammelost respectively:

Kürsteiner, J., and Staub, W. *Merklblätter für Emmentalerkäser*. Bern, Schweizer Käse-Union, A.G. etc., 1950.

Funder, S. *The Chief Moulds in Gammelost*. Oslo, Brøggers, 1946.

On inhibitory substances, including penicillin:

Mattick, A. T. R., and Hirsch, A. *International Dairy Congress*, **2**, 546 (1949).

Hunter, G. J. E. *Journal of Dairy Research*, **16**, 235 (1949).

Hirsch, A., McClintock, Mary, and Mocquot, G. *Journal of Dairy Research*, **19**, 179 (1952).

For some very recent work, see the reports of the 13th International Dairy Congress, The Hague, June 1953.

GERMINATION

A. M. MAYER

ONE of the most basic, and also one of the most intriguing processes of plant life is that of germination. We can define germination as that group of processes which causes the sudden transformation of the dry seed into the young seedling. This definition is very useful because it is what we normally see when watching germination occurring in the garden or in the field. In other words this is the germination which interests the farmer or the gardener. For them, germination is only of interest when something, the first leaves or the cotyledons, emerges from the ground.

To the plant physiologist, however, the process is much narrower, and at the same time less well defined. For him the emergence of some visible organ such as the young root or the young shoot is only the final manifestation of germination, not its beginning. It is only a sign that an extremely complex series of physiological processes has occurred. It does not and cannot tell us anything as to their nature, identity, or sequence of occurrence.

We know best the conditions required to bring about germination. Everyone knows that the first prerequisite of germination of the dry seed is an adequate supply of water. If other conditions are suitable, the first step of germination is that the seed imbibes water. This is at first a purely passive process conditioned only by the colloidal nature of the material of which the seed is built. The extent to which such imbibition will occur will depend primarily on the nature of the seed. Thus a rather starchy seed like the wheat grain will imbibe relatively little water, while a proteinaceous one like the pea will absorb a great deal. The rate and extent of water imbibition, irrespective of the seed, is largely dependent on the temperature. At high temperatures the absorption will be fast, at low temperatures it will be rather slow.

At the same time, however, at low temperatures the final amount of water which enters the seed will be rather larger than at high temperatures.

Equally characteristic of this primary step of germination is the increase in volume of the seed. These increases in volume go hand in hand with the development of large pressures. This can be simply illustrated by closely packing seeds, for example of peas, in a container, adding water to it, and then tightly closing it. A glass container will generally burst under these conditions. These changes in volume and consequent pressures may in part serve the germinating seed by eventually permitting the emergence of the delicate root and shoot from the seed coat without damage. It must, however, be stressed that they are quite inevitable physical processes. Dead seeds will absorb water for a considerable time just like live ones. The absorption of water, although it precedes germination, is by no means a proof of it.

The second requirement for germination is a suitable temperature. While, during storage, seeds show a very wide range of tolerance to temperature extremes, from well below 0°C . to 50°C . or even higher, this range narrows considerably during germination. For germination each variety of seed has its own optimum temperature and will germinate to any appreciable extent only within a very narrow region about this temperature. This optimal temperature can generally be related to the conditions of temperature under which the seed is likely to germinate under natural conditions.

Strangely enough, the nature of the effect of temperature on germination is only beginning to be understood. The amount of experimental data is prodigious. Thus we know today that the temperature response of seeds is conditioned by light. For example, certain lettuce seeds germinate very well in the dark at 26°C ., but fail to germinate at 29°C ., unless they first receive some light. Cases where the seeds will only germinate when kept in the dark and are adversely affected by light are equally well known. Wheat, for example, germinates best in the dark. A case where the relation between light and temperature is the reverse of that in lettuce is provided by timothy grass, in which tempera-

ture and darkness are interrelated. At one time it was thought that seeds could be divided into three classes: those germinating in the light, those germinating in the dark, and those indifferent to light. Recent work has disproved this view. Today we can say that all seeds will respond to light (or darkness), provided that the other conditions, and especially the temperature, are right.

Another essential for germination is an adequate supply of oxygen to the seeds. It is not the external oxygen concentration which is the determining factor here, but the amount of oxygen which can reach the embryo. The failure of many seeds to germinate can undoubtedly be ascribed to the fact that oxygen cannot penetrate into the seed, because, for example, the seed coat hinders it. The same is also true for penetration by water. Both can be improved greatly by slightly damaging the seed coat or softening it by chemical treatment, such as steeping the seeds for some time in acid. In the most extreme cases, where the seed coat is very hard and impermeable, we can cause germination by entirely removing it.

Thus far we have merely considered the external factors affecting germination. But internal factors are of equally great interest. The most obvious of these is the viability of the seed, that is to say, its ability to germinate. In effect this means merely whether we are dealing with a live seed or a dead one. This can nowadays be determined by staining the seed with tetrazolium chloride, a dye which gives a characteristic red colour with live seeds but not with dead ones. It is a useful method for sampling seeds of doubtful viability. Next are the conditions under which the seeds were produced. Every seed grower is acquainted with the very undesirable phenomena that seeds produced in certain localities, under certain conditions, fail to germinate, or germinate in less proportion than usual. Although this causes large annual losses to the seed grower and to the farmers who have to use the seed, we know practically nothing about its cause. Equally little understood is the periodicity which seeds show in their germination. Some seeds will germinate immediately after harvesting, but will then slowly enter a period of non-germination, called dormancy. After a period of weeks, months, or years

this so-called dormancy comes to an end, and the seeds again germinate normally. This period of rest can be artificially broken by treatment of the seeds with ethylene, thiourea, or alternating high and low temperatures. Other seeds will only begin to germinate after a certain period from harvesting has elapsed, this being generally termed after-ripening. But even more startling and inexplicable is the real periodicity which many seeds exhibit: under definite, constant conditions the seeds show periodic maxima and minima in their germination percentage. These maxima and minima may occur monthly or yearly, or after arbitrary but regular time intervals. As yet no one has succeeded in relating these periods to any definite process in the seeds. This is not unduly surprising if we now examine the sum total of our knowledge of the physiology and biochemistry of germination.

It is known that as the seed imbibes water it starts to respire at an increasing rate. This rate is by no means constant. Such an increase in respiration characterizes germination as opposed to passive water-uptake. An equally obvious feature of germination is the decomposition of some of the reserve materials present in the seed and their transformation into other substances such as cell material and new protoplasm on the one hand and their utilization for respiration on the other hand. These physiological processes are undoubtedly those around which the process of germination centres. At the start of germination we have an embryo which is more or less in a state of rest. After a very short period of time this embryo is suddenly stirred into activity and begins to develop. This development must in the first instance express itself in cell division, that is to say the formation of new cells. Consequently new cellular material must be created, whose source during germination can only be the material already present in the seed. The subsequent stages of development, cell elongation and vacuolation, do not strictly speaking belong to germination proper. It may be assumed that they are no different from the similar processes occurring in any growing tissue of the adult plant.

It is therefore the first step, the activation of the embryo,

which constitutes the real germination. This is evidently a very complex process. One of its first phases must be the hydration of the proteins of the seed, in order to make them potentially able to metabolize. But before the seed can metabolize, all the enzymes which are so essential a part of cell activity must become active. In the dry seed apparently they are not active, except those concerned with respiration, since even the dry seed respire, although at a very much reduced rate. It is then, when these enzymes have become active, that germination gets into full swing. At once there are a multitude of activities in the seed. The nuclei now begin to divide, reserve materials are hydrolysed to smaller entities and are reconstructed into new materials. Thus starch is broken down to glucose and fats are hydrolysed. The proteins are wholly or partially broken down and new and different ones are formed. Even such unpromising material as a date stone, which consists essentially of very hard cell walls, is utilized while the date embryo stirs into activity. New cell walls are laid down, and finally, the externally visible signs of germination begin to appear.

This part of germination is therefore of the greatest interest to the biologist. It provides us with material for studying the change over during a matter of hours of an almost inert, dehydrated mass of protoplasm to a highly active, organized, hydrated plant cell. The significance of this is self-evident. It can help us to understand the conversion of normally resting cells into active ones, as in the case of tumorous cells. Considering the significance of the process it is surprising how little effort has been devoted to its study and how great is our ignorance of it. The reasons for this are manifold. The most obvious is the enormous diversity of processes which are packed into a very short period of time. It is rather like trying to analyse what constitutes the really vital part of building a house in the absence of the blue print. At one moment there are just heaps of bricks and mortar lying around; at the next, dozens of workmen are beginning at many points – building walls, and putting in plumbing, electric wiring and so on. If we watch only the bricks, we shall be surprised at the end to find that the water taps work.

One very useful and interesting approach to the whole problem has, however, been made. It was suggested by the fact that in nature many seeds and fruits are provided with a regulatory system. This system consists of the presence in the seed of inhibitors which normally prevent their germination unless certain conditions are complied with. These conditions may be the correct climatic conditions such as an adequate amount of rain. This is a frequent feature in desert plants such as the guayule. Many fruits such as the tomato or the citrus also contain inhibitors of germination. Here the function may be to prevent germination prior to full development. The object of such inhibitory systems in nature, if one can refer to an object at all, is apparently to ensure that the seeds will only germinate under favourable external conditions. This safety mechanism provides an ideal tool for research. Evidently such an inhibitor must stop some basic process of germination. We can isolate such an inhibitor and study its effect on germination. Following this, we can investigate at what time in germination it acts, and can look for the process with which it interferes and identify the enzymes concerned. If we return to our analogy of the house for a moment, it is like forcibly preventing the plumber from working and then thoroughly investigating the difference in the house with and without the plumber working.

It is this method which has been chosen as one of the lines of approach to germination physiology in the laboratories of the Botany Department of the Hebrew University. It has been found that coumarin, a substance frequently occurring in seeds and fruits, is a very powerful inhibitor of germination. Its effect on various enzyme systems is being systematically studied. Moreover the several functional groups of coumarin have been investigated in the hope of elucidating the nature of its action. While it is too early to draw conclusions from this work, it is already clear that it is very promising and will eventually lead to a better appreciation of germination.

In research it is always easier to study the difference between two more or less similar processes than to investigate the nature of one of them without having a basis for comparison. In ger-

mination studies we make use of this, as well as of the general experience that the study of metabolism under abnormal conditions, as instanced by our inhibiting of germination, often provides a short cut to knowledge of the normal pathway of metabolism.

Germination is a process on which much of our agriculture and horticulture is based. It is frequent and obvious, and at the same time provides us with a most fascinating and fruitful field for future research. The results obtained are of interest not only because they concern germination, but because they may provide the key to the general problem of what constitutes activity or inactivity of protoplasm, the building material of living organisms.

FURTHER READING

For a general introduction, not including more recent work, see: Barton, Lela, V., and Crocker, William. *Twenty Years of Seed Research*. London, Faber, 1948.

PHENOMENA OF SUPERSONIC FLIGHT

G. N. LANCE

RECENT events have aroused public interest in two rather specialized phenomena of high speed flight. The first of these phenomena is that of 'Supersonic Bangs', which were first publicly demonstrated at the Society of British Aircraft Constructors' Flying Display at Farnborough in September 1952. The second, rather unexpected happening, is 'Control Reversal', which was brought into prominence in the recent film 'The Sound Barrier'. The object of the present article is to present an acceptable, if slightly simplified, explanation of the facts behind supersonic bangs and control reversal.*

* SUPERSONIC BANGS

How do subsonic, transsonic, and supersonic flight differ from one another, and why do they differ? The speed of sound in a fluid depends on the fluid itself and on its physical state. For example, if the atmosphere consisted solely of hydrogen, sound would travel very much more quickly than, in fact, it does in our atmosphere consisting primarily of nitrogen and oxygen. Supposing the atmosphere given, then the speed of sound depends on the height above sea level – it decreases with increasing height – also the speed of sound varies, slightly, from day to day as the barometric pressure alters. In fact it can be shown, in a rigorous mathematical way, that the speed of sound depends on the pressure, the density and the adiabatic constant of the propagating medium.

To understand clearly the difference between sub- and super-

*For a further account of aeronautics see the articles by Prof. O. G. Sutton and H. Pearson in *Science News* 7 and 18 respectively.

sonic speeds it is best to consider a simple example. Figure 1a shows a slowly flowing river into which pebbles are dropped; in Figure 1b the speed of the river is increased until it is exactly the same as that of the disturbances, and in Figure 1c the river is flowing very swiftly. The point O is the position of the observer who drops the pebbles, and is a fixed point in space; if the water is not flowing then the disturbances will consist simply of con-

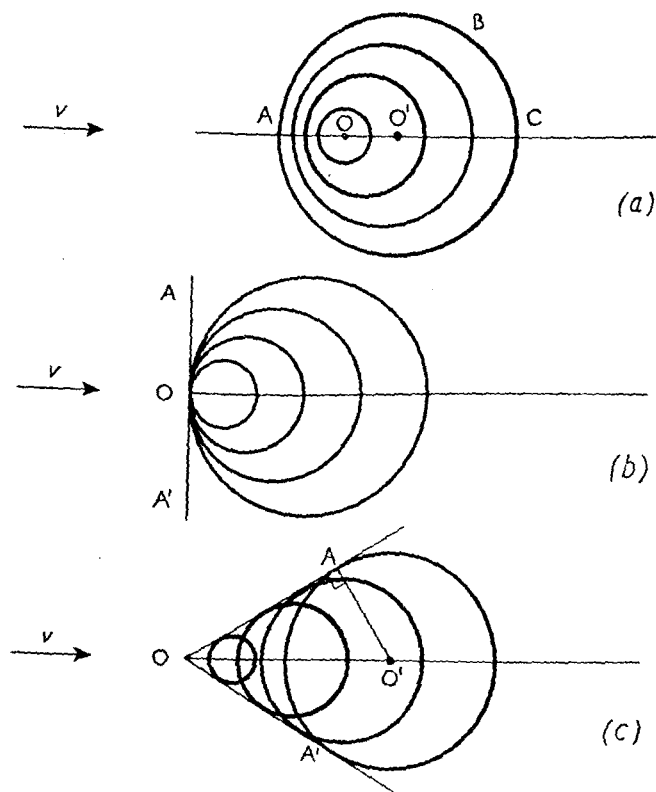


Fig. 1. The disturbance produced when pebbles are dropped into (a) a slow-moving stream, (b) stream flowing at exactly the same speed as the disturbance travels, (c) fast-moving stream.

Phenomena of Supersonic Flight

93

centric circles, their circumferences extend outwards from O with the speed of sound (in water), say a . Now consider Figure 1a, in which the stream is flowing from left to right with a velocity V and the circles of disturbance are no longer concentric. After a time t the water will have travelled a distance Vt ; i.e. $OO' = Vt$, where O' is the centre of the circle ABC caused by the pebble thrown in at O at time $t=0$. It can be seen that the disturbances are 'concentrated' upstream from O and 'rarefied' downstream. The important thing to note about this case – the slow-moving stream – is that the disturbances could, if they did not die out, reach all points in space upstream or downstream from O .

Turning now to Figure 1b, in which we suppose that $V=a$; then the point O will be such that the disturbances cannot move upstream against the current. Thus the disturbances form a system of circles with a common tangent $A'O A$ which is perpendicular to the direction of the stream. In this case ($V=a$) it can be seen that the disturbances can only reach points which are downstream from $A'O A$, and no point upstream from $A'O A$ can ever feel the effect of the pebbles.

Finally, consider Figure 1c in which V is greater than a . Here it can easily be proved that the circles of disturbance all touch two common tangents, OA and OA' . Since $AO' = at$, and $OO' = Vt$, it is clear that

$$\sin \angle AOO' = AO'/OO' = a/V = 1/M. \quad (\text{Eqn 1})$$

The non-dimensional quantity $M = V/a$ is defined as the 'Mach Number' of the flow, and the angle $AOO' = \mu$ is called the 'Mach Angle'. Here again, the disturbances cannot reach all points in the fluid but are confined to a wedge, whose angle 2μ depends on M . The connection between μ and M is given by Eqn 1; the greater the speed, V , the smaller is the angle, μ .

Single Bangs

With these preliminary ideas, it is now possible to return to the subject. It makes no difference to the argument whether the stream flows with a velocity V and the pebbles are dropped in at the same point in space or whether we drop the pebbles into still

water but do so when running with a speed V . Furthermore, we obtain a much closer analogy with aeroplane flight if we think of the disturbances in three dimensions instead of only two as we have up to the present. The extension is easy; circles become spheres, the common tangent $A'OA$ becomes a tangent plane, and the wedge $A'OA$ becomes a cone – the Mach Cone. The disturbances can be caused by rapidly inflating and deflating a succession of balloons arranged along the flight path. This analogy is the one used by C. H. E. Warren,* and will prove very helpful in understanding the bangs.

It should be pointed out that, if the velocity of flight does not remain constant, then the diagrams will be altered slightly, e.g. the lines OA and OA' of Figure 1c will no longer be straight.

What of the bangs themselves? The most accurate description of them is that they sound similar to one, two, or three rounds of anti-aircraft fire. It is rather difficult to define the difference between a bang and a loud noise, but a bang is best described as a very intense noise of short duration. Any explanation of the bangs must account for the following facts:

(1) As many as four bangs in rapid succession have been heard by observers.

(2) A number of observers listening to the same flight may not agree as to the number of bangs they heard; this is an important point.

(3) The bangs may be heard by observers scattered over an area of some 10 by 4 miles.

Anyone who has heard a jet-engine on the ground will agree that, if this noise could be stored up or accumulated in some way and then fired at an observer, this observer would hear a very loud noise. If the aircraft were capable of flying for a considerable period of time at approximately the speed of sound ($M=1$) then the acoustic noise would build up at the aeroplane as at O in Figure 1b. When the aircraft slowed down, this large disturbance would then continue to travel with the speed of sound, a , and any observer who was in its path would no doubt hear a

*See the correspondence in *Flight* quoted in 'Further Reading'.

loud noise. Unfortunately this explanation is rather too simple and it must be rejected for the following reason:

Unless the aircraft were capable of *sustained* flight at about $M=1$ the sound of its engines, etc., would not accumulate; and, in fact, the flight records do not show that the aircraft flies with a velocity of $V=a$ for a sufficiently long period of time to cause a build-up of sound as described above.

The possibility must now be considered of the aircraft creating large aerodynamic, as opposed to acoustic, disturbances. (Since a balloon emits no noise, the disturbances, discussed above, which are caused by hypothetical balloons are aerodynamic rather than acoustic.) Aircraft do, in fact, produce such disturbances and their occurrence may be explained best by the balloon analogy. In place of the aircraft we imagine that as it flies along a balloon is suddenly inflated and fills the space which the aircraft would occupy and then the balloon deflates again as the aircraft passes. These closely spaced balloons cause aerodynamic disturbances (or sudden changes in pressure and density) to propagate, with the speed of sound, as in Figures 1a, b, or c, according as the imaginary aircraft is flying sub-, trans-, or super-sonically.

It is now well known that certain aircraft are capable of sustained supersonic flight, and this is taken as a basis for our next explanation. Warren asserts that the supersonic bangs can be heard by any observer who is stationed at a point on the ground towards which the *component* of velocity of the aircraft is exactly a . Figure 2 shows the geometry of this condition.

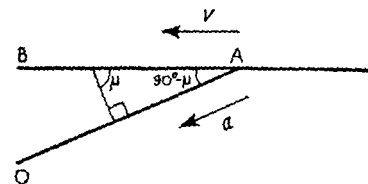


Fig. 2. Aircraft, A , flying supersonically with a velocity V will have a component of velocity equal to a in a direction AO , where $\angle BAO = 90^\circ - \mu$, and μ is the Mach angle.

Here AB is the path along which the aircraft A is flying supersonically; if $\widehat{BAO} = 90^\circ - \mu$, where μ is the Mach angle, then the component of velocity of the aircraft in the direction AO is a , the velocity of sound. Thus the aircraft has a component of velocity equal to a towards any observer who is stationed along the generator of a cone, vertex A , axis the flight path AB and semiangle equal to the complement of the Mach angle. Such an observer will, Warren proves, hear a bang. The present writer subscribes to this explanation as being, in his opinion, the best one so far advanced which embraces all the facts.

It can be shown mathematically that an observer so situated will receive the disturbances from many balloons simultaneously, very much as in Figure 1b, and that this disturbance is sufficient to cause a bang instead of a loud noise. Clearly, on this explanation many observers are all able to hear a bang, and the area over which it may be heard will be roughly rectangular, with the longer side in the direction of flight. Thus this theory can account for at least one of the facts listed above. It now remains to explain the facts (1) and (2): that more than one bang may be heard by some observers.

Multiple Bangs

A typical flight path is shown in Figure 3; P is the point at which the aircraft begins its dive, and, before reaching A_1 , it has

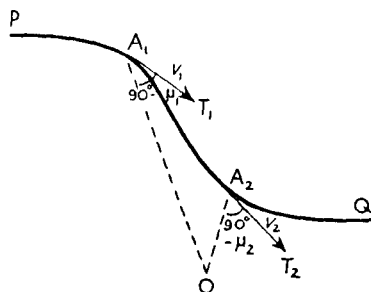


Fig. 3. At the point O two bangs will be heard; that produced at A_1 will arrive at O first and that produced at A_2 will arrive later.

attained supersonic speed. The plane continues to accelerate for a while, and at A_2 has a velocity V_2 greater than V_1 , that which it had at A_1 . Since V_2 is greater than V_1 , it follows that the Mach angles at A_1 and A_2 are such that μ_1 is greater than μ_2 ; and, if A_1T_1 and A_2T_2 are the tangents to the path at A_1 and A_2 respectively, then $\widehat{OA_1T_1} < \widehat{OA_2T_2}$. Hence it is possible for a point, such as O , to lie along the line whose direction gives a component of velocity equal to a twice during the same flight. If this is the case then the observer stationed at O will hear two bangs, and a little thought will make it apparent that the bang heard first will have been caused by the aircraft whilst at A_2 and that heard second will have been caused whilst at A_1 .

The reader can convince himself by drawing other diagrams of possible flight paths and assuming other velocity distributions along the flight path that it is possible to find points at which three lines intersect along all of which the component of velocity is a . Observers stationed at such intersections will hear

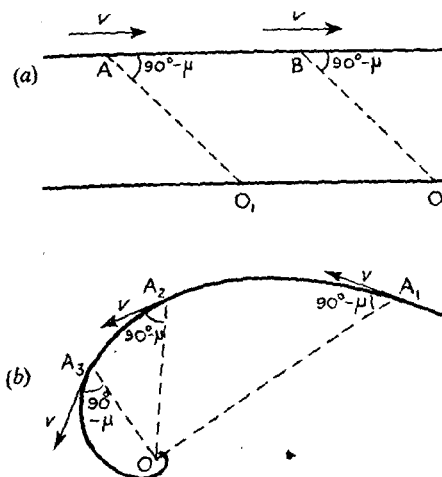


Fig. 4. (a) Straight flight path; single bangs will be heard along O_1, O_2 . (b) Equiangular spiral flight path. The focus O receives a bang from all points of the path. (See following page.)

three bangs. In this way we have been able to explain the other two facts about bangs; namely, that more than one bang may be heard by any given observer, and that the number of bangs heard by different observers is not necessarily the same.

Reports have been published of bangs being so violent that windows have been shattered by the force of the 'explosion'. If supersonic flights are to become everyday occurrences then it is important to investigate the possibility of the bangs being very violent. Consider for the moment a straight flight path with a constant supersonic velocity V parallel to the ground. (See Figure 4a.) In the diagram the aircraft will clearly cause bangs all along the ground and any point O on the ground will hear its bang from only one point, A , of the flight path.

Suppose now that the flight path were twisted in such a way that the observer O were to hear a bang from the plane wherever it was on its flight path – or at least for a considerable portion of it – then the bang would certainly build up and become a violent explosion. Such a state of affairs could pertain if the aircraft flew, with a constant velocity V (greater than a), along a portion of an equiangular spiral. (Figure 4b.) However, if the aircraft did not maintain a constant velocity then the exact shape of the flight path required to produce an explosion at a given point would not be such a simple curve, and in fact the path would be a function of the velocity at any point. Hence to prevent violent explosions pilots of the future must be instructed not to travel along equiangular spirals with constant supersonic velocity!

There is one other explanation, which was put forward by B. A. Hunn, and this will now be considered. Hunn believes that it is flight at approximately sonic speed which causes the bang and he also showed that it is the *aerodynamic* as opposed to *acoustic* disturbances which are built up. The fairly large area over which bangs may be heard is accounted for on this theory by the facts that (a) the bang 'spreads out' after having left the aircraft when the latter slows down, and (b) the temperature gradient in the atmosphere causes the direction of the bang to be altered continuously.

Hunn is able to explain multiple bangs by saying that a bang is produced each time the aircraft passes through or 'touches' the speed of sound; thus an aircraft with a flight history similar to that shown in Figure 5 would cause three bangs. Finally, the difference in the number of bangs heard by different observers is explained by again making use of the refracting properties of the atmosphere; thus one or more of the bangs may not reach all observers, because their directions are altered by different amounts since the bangs were generated initially at different heights.

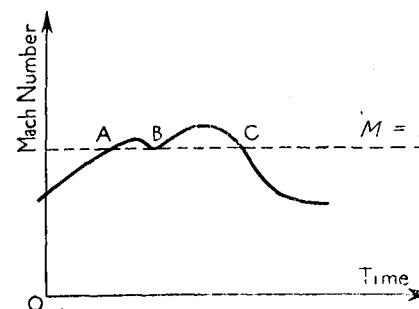


Fig. 5. Flight history which would, on Hunn's theory, produce three bangs when sonic velocity is attained at A , B , and C respectively.

To summarize the present position, the theories of Warren and Hunn each have something in their favour, and the correct one will only be decided when an aircraft is flown in steady level flight at about $M = 1.5$. If, as I believe, bangs are still heard then Warren's theory is more probable; if, however, bangs are not heard, then the true explanation clearly depends on *trans-* and *not* supersonic flight, and Hunn wins the day.

CONTROL REVERSAL

Before attempting to put forward an explanation of control reversal we must be clear as to what the expression means. For clarity we must confine ourselves to aileron reversal. The ailerons on an aircraft are usually situated on the trailing edges

of the wings and extend approximately halfway along the wing from the tips. They are hinged to the wing in such a way as to allow up-and-down movements of their trailing edges, and they are actuated by a sideways movement of the stick; if the stick is moved so as to cause a downward movement of the aileron on the starboard wing, then the port aileron will have an upward movement. In such a motion the starboard aileron is said to have a positive incidence and the other a negative incidence; and normally the starboard wing will rise. The rate at which the wing rises, causing the aircraft to roll, depends on several things, two of which are the amount of deflection of the aileron and the forward velocity of the aircraft.

It is found in practice that, for a given aileron deflection, the rate of roll becomes less as the forward speed of the aircraft becomes greater until a velocity is reached at which a deflection of the aileron produces no rolling at all. If the aircraft velocity is still further increased, then a positive incidence of the starboard aileron will cause the starboard wing to drop, instead of rise, as is normal with positive incidence, and the ailerons are said to be reversed. The critical speed at which no rolling occurs is called the 'Aileron Reversal Speed'. It may be pointed out here and now that the aileron reversal speed of an aircraft is *not necessarily the same as the speed of sound*, and in fact the concept of the speed of sound can from now on be forgotten. One of the aims of an aircraft designer is to make the reversal speed

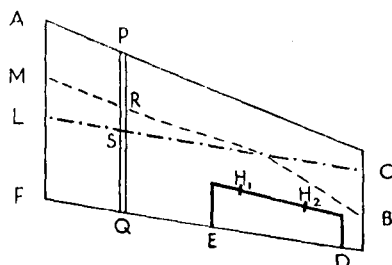


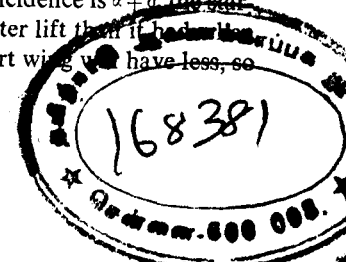
Fig. 6. Swept-back wing showing aileron attached by the hinges H_1 and H_2 . Centres of pressure locus is shown ---- and the flexural axis ———.

greater than the maximum speed at which the aircraft is designed to fly. In this way control reversal can and should be avoided.

Figure 6 shows a typical swept-back wing with the aileron not yet deflected, attached by the hinges H_1 and H_2 to the wing. Assume for the moment that the wing is perfectly rigid and will not bend under any forces that may be applied to it; it is inclined at an angle of incidence α . In this case the lift of a small section PQ may be considered to act at a centre of pressure R, and similarly for all possible elementary strips. In this way the dashed line MRB is obtained; it is called the centre of pressure locus. It is, of course, impossible at the present time to construct a perfectly rigid wing, and all wings must have certain elastic properties. Suppose a vertical force to act at different points in the strip PQ. There is one point only in that strip, shown as S, at which the force may act without causing any distortion of the wing. A force at any other point in PQ will cause a change of incidence; it will increase α if it acts forward of S, and it will decrease α if it acts aft of S. Similarly for all other elementary strips. The locus LSC is called the flexural axis of the wing; it can often be well approximated to by a straight line.

When the elastic properties are taken into account the wing will distort under the aerodynamic lift which must be produced to keep the plane in the air. If the wing root AF is at an angle of incidence α then owing to this distortion the remainder of the wing will be at a slightly different incidence. The aerodynamic loads on the strips and the distortions produced by them automatically adjust themselves so that, in level flight, the lift is equal to the weight of aircraft.

Up to this point the aileron has been taken as part of the wing, and in fact was at the same incidence to the air stream as the wing itself. We now turn to a consideration of the aileron as a separate unit. In the rigid wing case, if the starboard aileron is deflected an extra amount δ so that its incidence is $\alpha + \delta$, the starboard wing as a whole will have a greater lift than if had the aileron was undeflected, and the port wing will have less, so the aircraft will roll.



What happens when elasticity is taken into account? Since the aileron load is transferred to the wing by hinges H_1 and H_2 , which are usually aft of the flexural axis, the wing will distort, its incidence will decrease, and the effect of the ailerons may be neutralized or even reversed, in which case no rolling will occur. Since the amount of flexure of the wing increases as the aircraft's speed increases the ailerons become gradually less and less effective until reversal speed is reached, when the aircraft rolls the 'wrong way'.

As explained above, for safety's sake aircraft designers must so arrange things that control reversal does not take place within the speed range for which the aircraft is designed. How may this be done? A complete answer to this question cannot be given in an article such as this, but two ways of preventing, or at least delaying, reversal can be mentioned. A reference to Figure 6 will make it obvious that, for a given force on the hinge points H_1 and H_2 , the resulting distortion of the wing will be proportional to the distance of H_1 and H_2 from the flexural axis LSC . Since the size and position of the aileron is determined from other considerations this distance can be reduced only by moving the flexural axis towards the trailing edge. This will involve detailed alterations of the ribs, etc., composing the wings. The second way of delaying aileron reversal is simply to make the wing stiffer by thickening the skin or by other means; however, the difficulty here is that usually more material is required to produce a stiffer wing, and this all adds to the weight of the aircraft.

In conclusion, it should be said that the problem of aileron reversal is rather more intricate than has been stated. No account has been taken of the fact that the wing distortion depends also on the density of the atmosphere; thus, for a given speed, reversal occurs first at low altitudes. Nothing has been said about the reversal of other control surfaces, e.g. elevators. To discuss *this* subject adequately the distortion of the fuselage must be taken into account. However, one thing has been made clear: that if the aircraft were made inelastic, then distortions could not occur, and consequently reversal would not take place. But there is still a reservation which is necessary to this last state-

ment. If the wings are made thick, reversal may still occur because at high subsonic Mach numbers shock waves are induced on the wings by the deflection of the ailerons; this latter is a very complicated phenomenon and cannot be discussed here. Finally, it should be re-emphasized that control reversal has nothing whatever to do with that magic thing from which we started, 'The Speed of Sound'.

FURTHER READING

Supersonic bangs have been discussed at length by many authors, see for example:

Flight, issues dated 3 October, 17 October, 31 October, 21 November and 5 December 1952.

Nature, issue of 31 January 1953.

The detailed theory of Aileron Reversal has been given in *The Aeronautical Quarterly*, Volume 1, May 1949.

RESEARCH REPORT

A. W. HASLETT

WHISTLING ATMOSPHERICS

ALTHOUGH whistling atmospherics had already been discovered and described by H. Barkhausen in 1919 – that is, one year before the earliest regular broadcasting service – until a recent investigation by L. R. O. Storey at the Cavendish Laboratory, Cambridge, little systematic research had been done on them. Their characteristics are that they are of audio frequency (from about 0.5 to 10 kilocycles); that they are of long duration (from half a second to several seconds); and that there is a fall in pitch, or frequency, as the 'whistler' proceeds. The effect of Mr Storey's investigations is to suggest that they offer a new and promising means for the study of the most remote regions of the atmosphere surrounding the earth. His experiments were illustrated at a *Conversazione* held lately by the Royal Society, and have since been published.

When whistlers were discovered, and for some years afterwards, the main interest in atmospherics was as unwanted disturbances to communication. The immediate problem was to determine the origins of the more frequent and short-lived 'click' type of atmospheric, of about 1 to 2 milliseconds duration, and to correlate these with the distribution of thunderstorms and meteorological fronts. It was from this work that the 'Sferic' service of the Meteorological Office was developed as an aid to forecasting. More lately, atmospherics have been used increasingly in the study of radio propagation.

One of the earlier observations on whistlers was that they were often heard one to two seconds after a 'click' atmospheric. Since click atmospherics had already been connected with lightning strokes, it was a reasonable guess that whistlers, which sometimes followed the clicks, had a similar origin; and a theory,

based largely on this coincidence, was suggested by Barkhausen in 1930, and developed by T. L. Eckersley of the Marconi Wireless Telegraph Company in 1935.

It may be useful at this point to state the implications of any theory which takes this coincidence as a starting point. If a whistler lasting a second or so is considered to have originated in a lightning stroke lasting 1 to 2 milliseconds, and taking place another second or so earlier, then the interval between click and whistler, the duration of the whistler itself, and the change in frequency as it proceeds must all be regarded as secondary effects, determined by conditions during a surprisingly long interval of propagation. This interval is longer in many whistlers than the time taken by radio waves to travel (effectively under vacuum conditions) from the earth to the moon and back again. Secondly, since the time-interval between the arrival of the highest and lowest frequencies of a whistler is comparable with that between the click atmospheric, if any, and the beginning of the whistler, it follows that during much of their journey the wave-packets which compose the whistler must have been propagated under highly refractive conditions.

There was little enough evidence at the time. But on one point, Eckersley's theory led to a relationship which could be tested. It was a prediction from his theory that the square root of the frequency of a whistler should be inversely proportional to the time-interval from its origin. This relationship was confirmed in 1933 for a single whistler, and to that extent the primary assumption made in the theory was supported by evidence.

The backbone of Storey's research was the systematic recording of large numbers of whistlers, and the analysis in terms of frequency of all good specimens recorded. He used a straightforward receiver, restricted to frequencies between 0.4 and 10 kilocycles, and during one minute in every half-hour, day and night, the output from this was recorded automatically on a magnetic tape recorder. Routine observations consisted in a daily play-back of recordings, the counting of whistlers, and the analysis of good ones. During three months in 1951 the average number recorded was about 1.3 per minute. In addition there

was a series of special experiments in co-operation with the Sferic organization. The object of these was to determine the origin of click atmospherics followed by whistlers. During these experiments, all 'fixes' of atmospherics by the Sferic organization were followed immediately by the transmission by telephone, from Dunstable to Cambridge, of a short tone which was recorded at Cambridge on the same tape as the output from the atmospherics receiver.

It was a general observation from the Cambridge records that whistlers either came after a loud click or none at all, and this was confirmed by comparison with Sferic fixes. Clicks preceding loud whistlers came always from within 1,000 kilometres from Cambridge, and clicks originating at distances greater than 2,000 kilometres were seldom followed by whistlers. They were typical click atmospherics in that they came from storm areas, and their wave-forms, when examined by P. W. A. Bowe, were found also to be typical. There is thus good evidence that whistlers preceded by clicks originate in lightning strokes in the same part of the world, and that they return after an interval.

What of whistlers not preceded by clicks? Here the relationship between frequency and time, already mentioned, has provided a useful indication. This relationship can be expressed in the form that the time multiplied by the square root of the frequency is equal to a constant D , which represents the dispersion of the whistler. (If the drop in frequency is the same, the longer of two whistlers has the bigger dispersion.) Both the dispersion and the presumed time of origin of a whistler can be determined from its frequency at different times, whether or not the particular whistler is preceded by a click. But where a click and a whistler are associated, the time of the click and the time of origin of the whistler are the same, within the limits of accuracy of the method.

From their dispersions, single whistlers can be classified as 'long' or 'short'. There are variations in the dispersions found at different times; but at any one time there is a long group and a short group. Whistlers not preceded by clicks are short. Long whistlers, on the other hand, are consistently preceded by clicks,

and their dispersions are roughly twice as great. The simplest interpretation is that long whistlers have travelled twice as far as short ones. In addition, there are successions of whistlers. In some of these the dispersions of all whistlers are the same, and the intervals between them are then often irregular. The simplest explanation is that they have followed the same path from the same area, and originated in a multiple lightning-flash. But there are also trains of whistlers, in which the dispersion of each whistler is progressively greater than that of its predecessor. If the dispersion of a typical short whistler at the time is taken as unity, then the successive dispersions of a train of whistlers are either 1, 3, 5 . . . or 2, 4, 6 . . ., according as the first whistler is short or long. In either case, the *extra* dispersion of each whistler, compared with the one before it, is equal to two such units.

Storey's explanation can be most conveniently introduced at this point. It is that the unit of dispersion, as thus defined, is that introduced by a single journey between the area round Cambridge and a corresponding area in the southern hemisphere. A short whistler – not preceded by a click – is thus supposed to have originated in a lightning stroke on the far side of the equator, and to have made only one such journey. On the other hand, a long whistler (preceded by a click) has originated in the area round Cambridge, but has travelled across the equator and back again, probably by roughly the same route. Trains of whistlers are produced by waves which have been reflected backwards and forwards, possibly as many as ten or more times; each whistler of a train has thus, compared with its predecessor, made two extra journeys across the equator.

The interpretation that long and short whistlers originate on opposite sides of the equator is supported also by such figures as are available on seasonal variation. In February and March 1951 the average numbers of short whistlers recorded per minute at Cambridge were respectively 9 and 10 times as great as the average number of long whistlers; but in July the number of short whistlers was reduced to about one-fifth of the March figure, and the number of long whistlers was increased by 40 per cent. Short whistlers were still in the majority; they have

only one journey to make, compared with two. But the directions of seasonal variation were opposite for the two types, and in agreement with the variation of thunderstorm activity in the two hemispheres.

The theory developed by Eckersley has been extended by Storey, and is shown to provide an explanation of these facts. It is necessarily complicated since the medium of propagation is both dispersive and also has directional properties in virtue of the earth's magnetic field. Eckersley showed that waves of these very low frequencies should penetrate the ionosphere,* and that their velocities of propagation should increase with frequency, as is necessary for the explanation of the falling frequency of a whistler in terms of a dispersive theory. Storey shows further that, over a wide range of initial directions, the waves must be so refracted that the ray-directions correspond quite closely with the lines of force of the earth's magnetic field. Events up to the upper limits of the recognized ionosphere should be adequately predicted by the theory. Thereafter, given a sufficient density of electrons, the waves will continue to be propagated in the general direction of the magnetic lines of force. Above the magnetic equator, they should reach a height of about two earth radii, and should then return to earth by a path symmetrical with that followed during the first half of their journey.

The only obstacle to the theory is the relatively high electron density – of the order of some hundreds of electrons per cubic centimetre – which it is necessary to assume at these great heights. On the other hand, it is a fact of observation that the whistlers return to earth, and there is rather strong evidence that they cross the equator. It is suggested that the electron cloud may be derived, not from the earth's atmosphere, but from the

*In radio communication, the frequencies used are much greater, and waves that penetrate the ionosphere are those with frequencies greater than the critical escape frequency. For an introduction to some aspects of radio propagation, see *Fundamental Principles of Ionospheric Transmission* (Radio Research Special Report No. 17, H.M.S.O., 1948). Whistler propagation is taken to be in the 'extraordinary' mode (cf. doubly refracting crystals in the optical case).

sun. Storey points out in a note added to his paper in February that the zodiacal light is believed to be caused by the scattering of solar radiation from a disk-shaped cloud of tenuous material which surrounds the sun and is extended in the plane of the earth's orbit. From recent observations, H. Siedentopf, A. Behr, and H. Elsässer have estimated the density of free electrons at various distances from the sun, that at the distance of the earth's orbit being of the same order as is required by Storey's theory of whistler propagation.

Storey, L. R. O. *Philosophical Transactions of the Royal Society*, **A**, **246**, 113 (1953).

Siedentopf, H., Behr, A., and Elsässer, H. *Nature*, **171**, 1066 (1953).

BIOCHEMICAL COILS

In a note in *Science News* 22, it was recorded that a coiled (formally a helical) structure for protein had been put forward by Professor Linus Pauling and others, of the California Institute of Technology, and that evidence in its support had been published by M. F. Perutz of the Cavendish Laboratory, Cambridge, from x-ray reflexions. Events have followed a rather similar course in the case of another biologically important substance, deoxyribose nucleic acid. Its components were illustrated by D. L. Woodhouse and H. S. A. Sherratt in *Science News* 24 (pp. 29 and 33), and the evidence on its distribution, including that in chromosomes, was discussed. As with protein, several different types of coiled structures have been proposed. The earliest such suggestion, based on direct studies of nucleic acid fibres, appears to have been made by A. R. Stokes of the Wheatstone Physics Laboratory, King's College, London, and M. F. H. Wilkins of the Medical Research Council Biophysics Research Unit, also at King's College. The most recent and successful has been published by J. D. Watson and F. H. C. Crick of the Medical Research Council Unit for the Study of the Molecular Structure of Biological Systems at the Cavendish Laboratory, with supporting evidence from King's College.

The structure proposed for the salts of deoxyribose nucleic

acid consists of two chains coiled in the same direction about the same axis, but not equally spaced from each other along the axis. It could be imitated by winding two wires round a pencil, and then removing the pencil. The chains proper are necessarily formed of alternate phosphate and deoxyribose (sugar) components. Characteristic features are that the phosphate components are placed on the outside of the coiled chains, and that the bases, which are attached to the sugar components, extend inwards towards the axis, being weakly linked together by hydrogen-bonding, the bases of one coil to the bases of the other. If the axis of the fibre is imagined to be vertical, as in a staircase, then these cross-linkages are horizontal. There is no restriction from structure on the sequence of bases along any one coil, considered by itself. But given the sequence of the bases in one coil, there are restrictions on that in the other both from limitations of space and the possibility or otherwise of cross-linkages. With this structure, a purine base can be linked only to a pyrimidine base; or, more specifically, adenine (purine) to thymine (pyrimidine), or guanine (purine) to cytosine (pyrimidine). Hence, if the sequence of bases on one chain is considered to be determined, there is in each position on the other chain only one base* to which cross-linkage can be effected.

It follows that if cross-linkage is fully utilized, the ratios of the mutually linkable bases, adenine/thymine and guanine/cytosine, should be close to unity, although variation in the other two ratios, for example adenine/guanine, might be expected. Recent analytical results, quoted by Watson and Crick, show that for all sources of deoxyribose nucleic acid examined, these conditions are met.

The proposed structure offers several criteria for testing. It places the relatively massive phosphorus atoms at a distance of 10 angstrom units (10 Å.) from the axis; and accounts for previously known spacings of 3.4 Å. and about 34 Å. along the axis. The shorter of these distances is that between successive cross-

*Methyl cytosine, present in small amounts in deoxyribose nucleic acid from some tissues, is treated as equivalent to cytosine for this purpose.

linked bases, and the longer the interval at which the whole structure repeats. Lastly, there is the assumption of a two-coil arrangement, and the uneven spacing of the coils one from the other.

The evidence from x-ray reflexions, although not conclusive, provides either support or a measure of agreement on each point. Wilkins, Stokes, and H. R. Wilson discuss the reflexions given by the sodium salt of the nucleic acid from a variety of sources, including calf and pig thymus, wheat germ, herring sperm, human tissue, bacteria, and bacteriophage. The effect is to confirm the occurrence of a coiled structure of the dimensions stated, and also the placing of a large fraction of the total mass at the radial distance assigned to the phosphorus atoms. Evidence on the remaining two points is provided by a particular preparation from calf thymus, which has been examined by Rosalind E. Franklin and R. G. Gosling of the Wheatstone Physics Laboratory. 'The structural unit', they state from the evidence of x-ray reflexions, 'probably consists of two co-axial molecules which are not equally spaced along the fibre axis.' The latter point is critical to the whole theory of cross-linking by specifically matched bases, which is thus cautiously, and in part, confirmed.

Franklin and Gosling refer also to a different structure, also two-coiled, which corresponds with a crystalline form, and of which a detailed study is being published. It is formed in conditions of lower humidity, and is interconvertible with the other by the addition or removal of water.

With this amount of support, Watson and Crick have felt free to discuss the biological implications of their structure. They point out that if the two chains, initially cross-linked and coiled, became separated, each could form new cross-linkages only with the same bases, arranged in the same order, as in its original partner. Each would thus contain within itself the complete specification for the building up through new linkages of a complete two-coil arrangement. In this way, deoxyribose nucleic acids could be self-reproducing, even to the order of successive bases. And if the genes on chromosomes do indeed owe their

specificity to the arrangement of nucleic acid, as some have considered, both their general duplication and the occurrence of spontaneous mutation could be thus speculatively explained.

Watson, J. D., and Crick, F. H.C. *Nature*, **171**, 737 and 964 (1953).

Wilkins, M. H. F., Stokes, A. R., and Wilson, H. R. *Nature*, **171**, 738 (1953).

Franklin, Rosalind E., and Gosling, R. G. *Nature*, **171**, 740 (1953).

SCALE AND AGE OF THE UNIVERSE

Although much of the evidence has yet to be published, it is clear from preliminary accounts by Dr I. S. Bowen, Dr Walter Baade, and Dr Edwin Hubble that a major revision is in progress in the scale of extragalactic distances; and that, for interim purposes, the distances of all extragalactic nebulae can be treated as double their former values. Since the estimated 'age of the universe', based on the apparent rates of recession of extragalactic nebulae, involves estimates of their distances, this, too, will be doubled. The reason for revision is that the curve relating the absolute brightness and the period of variability of 'classical' cepheids, the variable stars used to provide a measure of distance, is now thought to have been in error. According to Dr Baade's account, the cause of trouble has been the treatment of two essentially different types of star as one type.

One of the earliest inconsistencies noted arose from an attempt by Dr Baade to find short-period 'cluster-type' variables in the central region of the Andromeda nebula. These are variable stars, with periods shorter than one day, which are named after the globular star clusters of our own galaxy. The periods of classical cepheids range from rather over one day to about fifty days. The cluster-type variables should have been visible by a small margin in plates exposed in the 200-inch telescope at Palomar Observatory; but they were not. Comparison with the brightest stars seen suggested a discrepancy of about $1\frac{1}{2}$ magnitudes, that is, of about four-fold in brightness. Further experiments were carried out to test the two obvious possibilities: that the cluster-type variables were fainter than had been supposed,

or that the distance of the Andromeda nebula was greater. To answer the first question, the 200-inch telescope was used to check the luminosities of cluster-type variables in the globular cluster Messier 3 by comparison with stars of 'solar brightness'. The results of these measurements suggested that the luminosities of the cluster-type variables were correctly known to within about one-quarter magnitude. To answer the second question, an independent estimate was made of the distance of the Andromeda nebula by photoelectric measurements made in three colours - yellow, blue, and ultra-violet - on some of the brightest stars in the nebula. From this information the luminosity and distance of a star can be calculated. The results agreed closely with the original estimate: namely, that there was an error equivalent to about $1\frac{1}{2}$ magnitudes in the distance of the Andromeda nebula.

Observations made on the Magellanic Clouds by Dr A. D. Thackeray and A. J. Wesselink have led to a similar conclusion. Distances estimated from cluster-type variables were greater by about three-quarters than the original estimates based on classical cepheid variables, and a preliminary account suggested that the final figures were likely to be greater if anything.

Dr Baade approached the problem of cepheid types through his discovery that star populations are of two types corresponding, as is now thought, with regions characterized by the presence or absence of interstellar matter. Stars of Population II, which are found in the latter case, include no young stars. In his own account: 'There were strong reasons to suspect that the pulsating cepheid variables, although present in both populations, really represented two different species and hence each group of cepheids had to be calibrated independently.' Cluster-type variables, with periods shorter than a day, exist only in Population II; in Dr Baade's nomenclature, they are 'Type II Cepheids'. But both types 'include stars whose repeated periods of brightening and dimming range from about two days to fifty'.

The probable revised estimates of the scale and age of the extragalactic universe seem likely to remove more difficulties than they create. One of these relates to the age of the earth,

which has been variously estimated at between 2,000 and 3,500 million years. Doubling the age of the universe – from 2,000 to 4,000 million years – should at least leave some margin. Secondly, it has been an old difficulty that our own galaxy has appeared to be considerably the largest of the known nebulae. Doubling the distances of the extragalactic nebulae doubles also their dimensions; and, if not wholly removed, this difficulty also will be reduced.

The revision of the scale of extragalactic distances now in progress may have a further value in relation to cosmology. It is a reminder, if one were needed, of the disproportion which exists between the observational foundations used by cosmologists and the theoretical structures based on them.

Journal of the British Astronomical Association, **63**, 133 (1953).

Thackeray, A. D., and Wesselink, A. J. *Nature*, **171**, 693.

California Institute of Technology, press release, 12 April 1953.

Hubble, E. *Monthly Notices of the Royal Astronomical Society* (in the press).

HIGH-ENERGY RADIATION AND PLASTICS

Changes produced in the physical properties of plastics by exposure to high-energy radiation in a nuclear reactor have been studied by A. Charlesby, and more lately by M. Ross, at the Atomic Energy Research Establishment, Harwell. The result in most plastics, including polythene and polystyrene, is that more cross-linking is introduced between the long chains of which the plastic consists. In physical properties, there is in polythene, first, a raising of the melting point; then, with greater irradiation, the polythene becomes more elastic, and above its usual melting point acquires rubber-like elasticity; and, finally, with enough irradiation, it becomes 'completely amorphous, hard, and glass-like in character', even at room temperature. The effect on cross-linking is proportional to the radiation received over a wide range of values, so that specimens with any required degree of cross-linking can be prepared. In the case of polymethyl methacrylate ('Perspex'), there is a different type of effect. There is a

loss of transparency accompanied by expansion, and a white, brittle material is formed which consists of a mass of bubbles. This effect is produced either by a relatively slight exposure to irradiation, followed by heating, or by a greater amount of irradiation without heating. The bubbling is attributed to the irradiation without heating. The bubbling is attributed to the breakdown of side-chains, leading to the release of carbon monoxide, carbon dioxide, and hydrogen; and the expansion which takes place may be to six or eight times the original volume. Similar effects are produced by gamma-radiation from cobalt-60 – equivalent to high-energy x rays – so that, if uses are found for this expanded plastic, there will be means available for its production without direct recourse to a nuclear reactor.

Charlesby, A. *Nature*, **171**, 167 (1953).

Charlesby, A., and Ross, M. *Nature*, **171**, 1153 (1953).

UNUSUAL INSTRUMENTS

Some instruments which made use of unusual principles were on view at a *Conversazione* held lately by the Royal Society in London. In one of them changes in the velocity of sound through a gas mixture were made to operate a warning system. It was designed by Dr L. E. Lawley and Dr. E. G. Richardson of King's College, Newcastle-upon-Tyne, in the first instance for use at a colliery where methane was extracted from boreholes near the workings and piped to the surface. Its purpose in this case was to give warning should leakage of air into the system result in a mixture approaching explosive proportions. A cathode-ray indicator shows the difference in phase between sound waves leaving a telephone earpiece at one end of a tube, through which the gas mixture passes, and a microphone at the other. When the proportion of one of the gases, and therefore the phase-shift, reaches a pre-set figure, it is arranged that a buzzer will sound and an external relay system can be operated.

An apparatus for the automatic analysis of coins and other archaeological specimens, without damage to them, was shown by Mr E. T. Hall of the Clarendon Laboratory, Oxford University. The specimen is irradiated by 'white' x-rays, and emits

characteristic secondary x-rays which fall on a rocksalt or calcite crystal mounted on an automatically scanning and counting spectrometer. The sensitivity is about 0.2 per cent, and, if it is known in advance which elements are to be watched for, the instrument can be pre-set so that only relevant wavelengths are scanned. Thus employed, it was suggested that it might also be of use in metallurgy where knowledge of the composition of a melt was required in a minute or two.

An instrument with a different type of interest has been designed by the Hydraulics Research Station, Wallingford, and was shown at the Physical Society's exhibition. Its purpose is to meet the difficulty that the miniature waves produced in flumes and basins for the study of coast and harbour problems are seldom of uniform height, whether because of reflexions in the model or of imperfections in the generator. The instrument is designed to count a sample of 300 waves, and to measure the height of each from trough to succeeding crest. A series of 34 vertical wires, which vary progressively in length, explore the level of the water. The shortest and longest are arranged to operate an alarm circuit if any wave is produced which is beyond the capacity of the apparatus. Between each of the remaining 32 probes and its neighbour a resistance is connected. As the water reaches successive probes additional sections of resistance are shorted out. A pulse-counting circuit counts the number of probes which are covered between trough and crest, and the effective height of the waves is calculated from message-register readings.

CORRESPONDENCE

SCIENCE AND SCIENTIFIC METHOD

MANY philosophers and a large number of physicists have attempted to describe science and scientific epistemology, but chemists and other scientists have rarely ventured into this field. Science itself is very difficult, perhaps impossible, of definition, and the best we can do may be to accept Polanyi's suggestion* that a scientist is one holding certain beliefs. These beliefs are concerned with the nature of problems and the ways in which they may be solved.

The most fundamental scientific belief is that there is order in the world, that *A* must necessarily be accompanied or followed by *B*. *A* and *B* need not be events — one may be a probability — but even the most enthusiastic opponent of determinism expects statistical predictions to lead to almost complete certainty on the macro-scale. This belief may or may not apply to all *A* and all *B*, some would apply it only to events in which we or other sentient animals consider ourselves to be observers, not affected by and not affecting the event under observation. In the field of the older sciences, physics, astronomy, etc., this leads to no difficulty, but in the sciences concerned with living beings it may be hard to decide whether the scientist is an observer in the above sense or not. Few would deny that a physiologist studying the circulation of the blood is, but what is the position of a psychiatrist? Perhaps such subjects will have to be recognized as sciences in which another observer describes the behaviour of an observer and a subject of observation.

This relationship of *A* to *B* is usually said to be governed by a law: a law thus states that if we observe *A* we shall then or later observe *B*. A law cannot say why this should be so; a law cannot

*Polanyi, M. *British Journal for the Philosophy of Science*, 3, 217 (1952).

tell us why water should boil at 100° C. under ordinary pressure, although it may state the problem in another form. It must be stressed that even on non-deterministic theories it is asserted that if one observes a sufficient number of *As*, one will also observe a certain number of *Bs*.

This belief in determinism is based on universal experience; the opinion has often been expressed that scientific laws are merely statements of past experience, that we should express no surprise if our house were to uproot itself, but this view is opposed to common sense, and certainly has no place in the scientific outlook. Indeed, if it were to occur science would immediately look for some law to explain its behaviour, and in the event of not finding one would invent one.

The second important belief of science is that within the whole field of phenomena there is only one explanation which will fit all the facts. All our knowledge (or, at least, that part which is the concern of science) reaches us through the senses, and science, true to its first belief, attempts to construct a 'picture' of something responsible for our sensations. Whether this picture is 'really' true is outside the scope of science, but science believes that only one picture will give a perfectly correct account of what will happen when we make all possible observations. This explains the eagerness with which science attempts to unify its laws, to replace a large number of specific laws by a few general laws. To use an analogy, a few square inches of red, blue, white, and so on, paint, are arranged so as to form a picture of a man or other subject.

Since the concept of order implies measurement, science is, as pointed out many years ago by Eddington, a study of measurements only. This is easily seen in physics, but even in chemistry the same is true. The identity of two compounds is assumed if they have the same properties, such as colour, melting point, density, viscosity, etc., all ascertainable by measurement. The statement that oxidation of acetaldehyde gives acetic acid means, in fact, that if we take a substance agreeing in all its measurable properties with a substance once called acetaldehyde and treat it with a substance similarly identical with a named

oxidizing agent we shall obtain a third substance measurably identical with the substance originally called acetic acid. As confirmation of this, it is well known that the failure of the reaction is usually attributed to impurity, i.e. it is asserted that the acetaldehyde or oxidizing agent was not identical in all its properties with the substance first so-called. The failure may also be blamed on wrong concentrations, temperature, etc., all measurable.

This example also illustrates the picture-making tendency of science: a picture is formed of atoms (themselves constituents of the picture) linked up (measurement again) in a definite way and under conditions which may be specified by measurement, linking up in another way. The tendency to unification is illustrated by the way in which this picture is used to explain other properties of acetaldehyde and eventually, the classification of all substances alike in certain ways as aldehydes, assumed to contain the same group of atoms.

It is thus necessary to specify the exact conditions for our measurements, and science, as is well known, has often to re-define common words* in the direction of greater clarity and rigor. The great reliance placed on sight is also explained: most measuring instruments are designed to be read by the eye.

In its attempts at unification science derives broad concepts such as energy, force, motion, and attempts to make the constituents of its picture as few as possible. I do not think that this can be justified if we are attempting to describe an external universe, but if we admit that the only goal of science is to predict the outcome of initial events, then it has obvious advantages.

On this view, all laws are to be regarded as successive approximations, but, when a part of the picture has worked well for many years, it will not be rejected unless its deficiencies are very well shown up. Thus the law of conservation of matter has served so well for so long that the recent objections to it will be required to show that they are able to take the whole, and more, of its place in the picture before they are accepted as scientific truths, i.e. parts of the scientific picture.

Practical scientists usually assume that the picture is indeed a

true representation of the world, and while this cannot be proved it seems idle to speculate on the nature of things for ever unknown to us. If there is an external universe its metrical attributes present themselves to us in the way described by science, so by our first scientific principle we assume that nature does not play jokes on us, and that she really is as described.

Tottington, Nr Bury, Lancs.

R. SHILTON

CONSCIOUSNESS AND THE BRAIN

The points raised in Mrs Knight's article 'Consciousness and the Brain' (*Science News* 25) and the subsequent correspondence in *Science News* 27 are of such vital importance to our whole outlook that perhaps you will allow some further discussion of this topic.

Dr Mays says that our immediate awareness of anything – say the sensation of seeing something red – and the scientific description of the attendant phenomena are 'different logical types'. I think this point needs taking much further. We may know that a red object emits radiation with a certain frequency spectrum, that this impinges on certain cells in our retina, and that these send out signals to other cells in the brain. The processes concerned must end somewhere, say in electro-chemical changes in nerve fibres or synapses, or whatever else future investigation may reveal. The real point is that such an event as an electro-chemical change is nothing but an electro-chemical change; it is not the sensation itself, and the same goes for any pattern of observable physical events. The sensation is unanalysable, undescribable, and unobservable by anyone except the subject. We may react to it and say 'I see something red', but that can have meaning to another person only because he interprets it by analogy with his own sensations. The same goes for other contents of our consciousness, such as thoughts and emotions.

A machine could be constructed which said 'I see something red' when a red object was placed in front of it, but this would

prove nothing except to him who confuses reaction to some stimulus with direct awareness of it, which is something quite distinct. Consciousness or direct awareness cannot be 'scientifically' observed in any other person, and our only knowledge of it comes from introspection.

We are conscious not only of each separate sensation but of the time relation of succeeding ones. We are also conscious of 'identity', i.e. of being the same self today as yesterday. These facts point to a stable entity as the background of our awareness, and it is this entity which we call 'mind'.

The above line of argument then leads to the conclusion that mind is something quite distinct from brain, although there is obviously some connexion between them, about which, to be frank, we know nothing. We may further conclude that the workings of 'mind' are essentially unexplainable in terms of physical events, and are therefore outside the realm of scientific method as this is understood today. Even our language may be inadequate for such a study. Perhaps the most promising line of investigation – if investigation is at all possible – is some form of introspection, using a discipline yet to be developed.

Certain tentative conclusions may perhaps be drawn at this stage of the argument. Our sensations are real if anything is (including even hallucinations – it is only their interpretation that is at fault), since all other notions of 'reality' are based on them. Hence mind *is*. When we say that something is or exists, we usually take it for granted that it exists somewhere in space and at some given time. It is, however, by no means certain that such categories are relevant when speaking of the existence of mind; indeed all indications seem to point to the contrary, since saying that something exists somewhere in space is almost tantamount to saying that it is observable. It might be necessary to assert that mind exists, but that it is nowhere in space, and perhaps does not exist in time either.

Such statements may be repugnant to the scientific mind, which is apt to dismiss them as metaphysical or mystical nonsense. If, however, science is now undertaking to investigate mind, it will have to think correctly about the subject, and either

accept the necessary inferences or give up the attempt. Of course there is no *a priori* certainty that mind is capable of understanding itself in the same sense that it is capable of understanding nature, and it may be that we shall have to be content with knowing that mind exists, and leave it at that.

London S.W.7

R. HADEKEL

COSMIC RADIATION

In the article on 'Cosmic Radiation' in *Science News* 27, mention was made of a possible diurnal variation in the large air showers observed by Elliot at Manchester, although it was also pointed out that the number of events was too small for an explanation in terms of statistical fluctuations to be ruled out. I am informed by Dr Elliot that subsequent observations have failed to confirm the effect, and that it now looks as if the explanation in terms of statistical fluctuation is the correct one.

E. P. GEORGE

Department of Physics, Sydney University

SOME BOOKS RECEIVED

SCIENTIFIC EXPLANATION: A STUDY OF THE FUNCTION OF THEORY, PROBABILITY AND LAW IN SCIENCE by R. B. Braithwaite (Cambridge University Press), pp. 376, 40s.

Professor Braithwaite's important but difficult book is developed from his Turner Lectures 1946. 'The business of a philosopher of science', he holds, 'is primarily to make clear what is happening in scientific thinking'; and it should be worth the time of most scientists to read him. An analysis of deductive systems and calculi leads to the conclusion that the preferred system is one in which all the difficult deductions are in the stage of pure mathematics, and applications to a particular subject-matter are left as late and made as easy as possible. The importance of freedom of theoretical terms is logically justified, and models, not surprisingly, receive scant favour. The discussion of probability extends to the criteria applicable to statistical laws, and presents the progress of science itself as a game against nature in which the theory of games may be helpful. The final section on law, which the theory of games may be helpful. The final section on law, taken to be empirical, seems to be directed more at philosophers than at scientists. Professor Braithwaite concludes that 'if clarification has the secondary effect of encouraging scientists to construct deductive systems and to use theoretical concepts freely ... it will assist in the progress of science ...'

ADVENTURES IN PHYSIOLOGY by Sir Henry Hallett Dale (Pergamon Press, 1953), pp. 652, £6 6s.

The publication of a selection from the published papers of Sir Henry Dale has been given historical value by the addition of an introduction and commentaries which supply the personal background to the work and the author's own judgement as to why particular experiments were done or not done, or conclusions drawn or not drawn, at particular times. The first impression is of the large part allotted to luck. Examples are the initial suggestion of a subject, ergot, which led to his two main interests, the transmission of impulses from nerve fibre endings and the study of histamine and shock, and two cases of errors in procedure turned strikingly to advantage. The second impression is that success in research depends, not on the

receipt of luck, but on the use made of it. And the third, valuable to all research workers, and several times illustrated, is the difficulty of reaching a new idea before its proper time. As its publishers and price should suggest, the volume is magnificently produced.

MR TOMPKINS LEARNS THE FACTS OF LIFE by G. Gamow (Cambridge University Press), pp. 88, 12s 6d.

Mr Tompkins's adventures in physics have been so pleasingly told and illustrated in earlier books by Professor Gamow that his present excursions into blood-stream and digestion, chromosomes and cytoplasm, and a brain populated by barking, Pavlovian dogs, hold an obvious and realized prospect of mixed entertainment and instruction. There is the usual concluding lecture, this time on 'the relation between the phenomena of life and the laws of physics'. Also, as usual, there are unobtrusive allusions which those who wish may seek to identify.

THE TOOLS OF SOCIAL SCIENCE by John Madge (Longmans, 1953), pp. 308, 25s.

The greater part of this book is an analytical account, with examples, references and frequent appraisal, of the methods of social science. In an introduction and long first chapter there are more general discussions on such matters as objectivity, presupposition, values, language, and the place of logic in research. The author's approach is pragmatic and unpretentious, and the more detailed chapters are never allowed for long to remain merely a catalogue. In a final discussion he claims a double position for social science – both as a science in the making and in society.

A BACKGROUND FOR THE PREVENTION OF BILHARZIA by Alan Mozley (H. K. Lewis, 1953), pp. 71, 9s.

The relationship between bilharzia, the parasites which cause the disease, and the snails which are their alternate hosts, has been used so often as an illustration in the teaching of biology that there should be a wide interest in this simple monograph, written primarily for public health workers. Dr Mozley gives most of his space to the characteristics, habitats, and limiting requirements of the snail species concerned – there are seldom more than two in an area – and for long-term control places the main emphasis on applied ecology; with

copper sulphate, knowledgeably used, as an immediately available measure. He predicts that more extended irrigation, particularly when perennial, will add to bilharzia areas.

IMMUNOCHEMISTRY: BIOCHEMICAL SOCIETY SYMPOSIA No. 10 organized and edited by Dr R. T. Williams (Cambridge University Press, 1953), pp. 94, 12s 6d.

The usefulness of this symposium is that it maintains a balance between the basic problems of immunology (mechanisms of immunity reactions and hypersensitivity), on which most commonly made statements are more simple and definite than the evidence warrants, and various current techniques (including radio-isotopic labelling, and the use of enzymes as antigens), which have not been often reviewed. It should be a useful supplement to the reading of medical students and others.

SOME ASPECTS OF ENZYME RESEARCH (*British Medical Bulletin*, Vol. 9, No. 2: London, British Council, 1953), pp. 84, 15s.

A symposium, too wide in scope to be readily summarized, arranged by a committee under the chairmanship of Sir Rudolph Peters.

WEATHER INFERENCE by D. J. Holland (Cambridge University Press), pp. 196, 30s.

This is a book without category. It is an attempt by a young professional meteorologist, evidently with Air Force experience, to provide an intelligent introduction to meteorology and weather forecasting through an expanded commentary on the weather that he observed and recorded in England during the autumn of 1936, when he was a boy without special facilities. A type of 'beginner' to whom it would be well suited would be a young biologist, lacking in thermodynamics, but amenable to the simpler physical formulae.

THE EARTH AND ITS MYSTERIES by G. W. Tyrrell (G. Bell and Sons, 1953), pp. 278, 16s.

Dr Tyrrell, formerly Lecturer in Geology in the University of Glasgow, and now an honorary research fellow of the same uni-

versity, has written a careful and well-illustrated book which is popular only in the good sense. It is the third member of a series which now covers the sciences of earth, sea, and sky, with the last three words of the title common to all books.

BRITISH PYRALID AND PLUME MOTHS by Bryan P. Beirne (Frederick Warne, 1952), pp. 208, 189 figs, 16 plates in colour, 21s.

It is forty years since the publication of a book in which all British species of the superfamily Pyraloidea are illustrated and their life-histories described. It is illuminating that of two species suggested for comparison with others as well studied, one should be *Ephesia elutella*, which as the 'Cacao Moth' or 'Tobacco Moth' is of quite evident importance.

PHYSICS AS A CAREER by Norman Clarke (Institute of Physics, 1952), pp. 70, 6s.

Practical information and advice for the would-be physicist.

ABOUT OUR CONTRIBUTORS

R. S. BARNES graduated at Manchester University and, during the war, worked at the Admiralty Signal Establishment on naval radar. He is now employed at the Atomic Energy Research Establishment, Harwell, in the Metallurgy Division, where he is investigating the diffusion of metals in metals.

JAMES BARNETT is a zoologist who did operational research for Coastal Command and the Admiralty during the war, and for the past two years has been working at the National Institute for Research in Dairying.

T. F. GASKELL graduated from Cambridge University and obtained his Ph.D. degree early in the war for seismic work on land and at sea. During the war he worked with the Royal Naval Scientific Service and Combined Operations. Subsequently, he joined the Anglo-Iranian Oil Company, and was in Iran until 1950. He then obtained leave of absence to take part as chief scientist in the world cruise of H.M.S. *Challenger*. This cruise was completed in the autumn of 1952.

F. H. GEORGE read both the Moral Science Tripos and the Mathematical Tripos at Cambridge, and also studied physiology in the Natural Science Tripos. For the last four years he has been lecturer in the Department of Psychology at Bristol University. This summer he is to visit Princeton University as a 'visiting professor' for one year. Apart from his special interests in psychology, philosophy and mathematics, he is also addicted to literature, rugby football, and cricket.

M. N. HILL was born in 1919. He graduated from Cambridge University in 1946 after six wartime years spent with the Royal Naval Scientific Service working on underwater counter-measures. Since 1946 he has been with the Department of Geodesy and Geophysics at Cambridge, and is a Fellow of King's College. His geophysical work has been chiefly concerned with seismic prospecting at sea.

G. N. LANCE was educated at Reigate Grammar School and King's College, London. He graduated with first-class honours in mathematics in 1949, obtained his M.Sc. in 1950, and in 1952 was awarded a Ph.D. for a thesis on 'Some Problems of Supersonic Aero-

dynamics'. In the same year he joined Hawker Aircraft, Ltd, as a mathematician in the Design Department. Hobbies: philately, carpentry, and gardening.

A. M. MAYER, aged 26, is a graduate of University College, London, who turned from chemistry (B.Sc.) to plant physiology (Ph.D.). Since 1950 he has worked in the Botany Department, Hebrew University, Jerusalem, teaching plant physiology, and doing research on germination and germination-inhibitors, and on the mass culture of *Chlorella* as a possible source of green fodder.

