

The different faces of time in physics

Harvey R. Brown

Faculty of Philosophy, University of Oxford, Radcliffe Observatory Quarter 555,
Woodstock Road, Oxford OX2 6GG, UK

E-mail: harvey.brown@philosophy.ox.ac.uk

Abstract. Time in physics has several faces and in an important sense, is not really a thing at all. This paper examines a number of the features of time in physics with particular emphasis on the role of clocks in relativistic physics.

1. What is time?

Here are two competing slogans related to the nature of time, the first of which has been attributed to both Albert Einstein and John Wheeler, while the second is mentioned by Richard Feynman: “*Time is what stops everything happening at once.*” and “*Time is what happens when nothing else happens.*” Behind the first slogan are arguably the notions of *multiplicity* or *change*. The basic idea is that the world presents itself in different states (a multiplicity), which we disentangle by associating them with distinct instants of time (hence change over time). It is as if time is Nature’s way to get around the logicians’ Principle of Non-Contradiction! All of this is a good first approximation to the fundamental nature of the beast: no change, no time. But this seems to suggest that time is parasitic on the multiplicity of things or states of affairs, that it has no independent reality – quite at odds with the second slogan.

Isaac Newton regarded space and time as “absolute” (a stance which Einstein is supposed to have refuted, as will be seen). But Newton was quite clear that space is not a substance because in his view it neither acts on nor is acted upon by other substances. He almost certainly regarded time in the same way: an entity, yes, but unlike any other. Four centuries later Richard Feynman, who in being rightly dismissive of the second slogan above, was a little more ‘operational’ in his thinking: “Maybe it is just as well if we face the fact that time is one of the things we probably cannot define (in the dictionary sense) ... What really matters anyway is not how we *define* time, but how we measure it.” ([1], section 5-1.). The “metric” nature of time will shortly be discussed, the fact that we can measure duration between events by way of ideal clocks. Before the proper definition of clocks is considered, it is worth quickly comparing a generic clock with a typical measuring device: a thermometer. One might imagine that a clock is a kind of thermometer of time.

So let us think of the old-fashioned mercury thermometer. Being held in the armpit or under the tongue, it measures body temperature by way of exchange of energy between the mercury column and the body. They act on each other. But what mutual action is taking place when a clock measures time? An inertial ideal clock (including its source of power) neither loses nor gains energy in doing its job: it “measures” but does not *feel* time because there is nothing to feel. As Newton might have argued, time is not a substance. Recognising this simple point arguably makes some of the strange phenomena that will be discussed below a little less baffling.

2. The arrow of time

A birthday party balloon spontaneously deflates if its end is not squeezed, but disappointingly for parents and grandparents the reverse process of spontaneous inflation is never seen. This is but one example of the myriad irreversible processes occurring everywhere around us, but it seems that it was only in the late 19th century that physicists started to appreciate the oddity of this state of affairs. For how is it to be reconciled with the fact that the fundamental microscopic laws of physics are *time reversal invariant*, which is to say that they do not pick out a privileged direction of time? If one inverts the sign of the temporal parameter in the equations – along possibly with other transformations of variables – they end up taking the same form. (Today it is known that there is a notable exception involving the physics of the weak nuclear interaction, but no-one thinks that this is the reason for the macroscopic arrow of time.) The puzzle deepens considerably when one reflects on the fact that irreversible processes we observe taking place across the Universe appear to evolve in the same temporal direction.

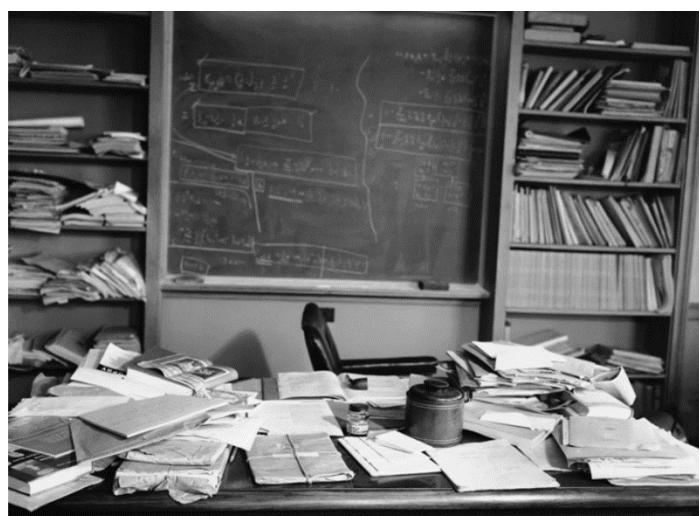


Figure 1. Einstein's office. "*The instant is not in time. Time is in the instant.*" (Barbour, [2])

There is a famous photo of Einstein's office in the Princeton Institute for Advanced Studies as it was when he died in 1955 (see figure 1). There are papers and books scattered across his desk, as well as bookcases groaning with collections of scientific journals. The photo contains information about what Einstein wrote and read in the past. While capturing only a single instant, it is a *time capsule*, a portal to history. So are fossils and sedimentary rock strata, which led geologists to discover 'deep time' in the history of our planet. The existence of such time capsules has led to the striking slogan due to Julian Barbour: "*The instant is not in time, time is in the instant.*" ([2], p 53). Other (if imperfect) portals to the past are our instantaneous memories. Why do we remember the past and not the future? Or better: why do our present memories give us information of varying degrees of reliability about events on one side of the present instant, which we happen to call "the past", and not events on the other side, which we happen to call "the future"? It is a banality to say history refers to the past, but it is anything but banal from the point of view of physics that there do not seem to be time capsules which are portals to the future. And of the portals that do exist, a remarkable feature of these, as Barbour has stressed, "*is the cornucopia of records that, hanging together with great mutual consistency, tell a story of a Universe of ever-increasing complexity and structural richness*" ([3], p 19).

The second law of thermodynamics – the increase or rather non-decrease of entropy associated with spontaneous processes in thermally isolated systems – is widely thought to capture the irreversibility of the relevant macroscopic processes. But the thermodynamic arrow of time is baked into the theory at an even deeper level. A more basic, often implicit, law states that systems out of equilibrium will spontaneously approach a unique state of equilibrium. Attempts that have been made using even more

fundamental microscopic physics to account for this irreversible equilibration postulate, as well as the law of entropy, almost always end up concluding that it is the special initial conditions of the system in question that provide the explanation. This situation is not unique. After all, when one rotates in one's chair the objects in one's line of sight can change and yet the fundamental laws of physics are (without exception) isotropic – they do not pick out a privileged direction in space. Ultimately the explanation here is tied up with the initial conditions of the Universe and perhaps the quantum fluctuations therein. It would seem the same goes for the thermodynamic arrow of time, and others (such as the electromagnetic arrow).

This attribution of very special initial conditions to the Universe is sometimes referred to as the “Past Hypothesis”. It may save the appearances, but it does not leave all physicists satisfied. Some hope that future physical laws will incorporate time asymmetry (e.g. [4], pp xi, 102). There has been much debate about whether the theory of the multiverse predicted by eternal inflation in cosmology can assign a probability to such initial conditions (see, e.g. [5], chapter 11). A more recent and perhaps more promising approach [3], [6] exploits the example of the time symmetric laws of Newtonian gravity to argue that there is a unique point in time in which the “complexity” of the Universe is minimised and then increases in both temporal directions for ever. It is consistent with the laws of thermodynamics holding with high probability within appropriate subsystems of the Universe in both directions away from this “Janus point”. Notably this approach goes some way to explaining the evolution of highly structured time-capsules in such subsystems. Time will tell which of these approaches, if any, is the right one.

Here follows a final word about phenomenological thermodynamics. The laws of this theory do not contain derivatives with respect to time. The theory of course deals with systems changing over time, but the laws do not say how fast they change: they do not refer to a temporal metric but merely a temporal order. In this sense thermodynamics contains no clocks and so is unique in physics (although one might argue that Einstein's 1915 field equations for gravity are likewise clockless until it is specified how the non-gravitational interactions couple to the metric (gravity) field.) Now let us turn to the case of theories which do contain temporal derivatives, such as classical electrodynamics – the theory of the interplay between electric and magnetic fields as well as electric charges as originally formulated by James Clerk Maxwell in the 1860s.

3. Duration: the metric of time

One of the great contributors to the theory of electrodynamics was the Irish physicist George Frances FitzGerald, who famously had the first premonition of relativistic length contraction (but not time dilation). At the end of the nineteenth century, FitzGerald was concerned with the definition of the standard of time in physics to be used in testing theories like mechanics and electrodynamics. He pointed out that the most accurate clock at the time was that used by astronomers: the rotation of the Earth. But he was aware that this leads to a conundrum. There are various reasons – for instance, the recorded times of ancient eclipses and the expected effects of tidal friction, which involves the oceans, atmosphere and the solid Earth itself – for thinking that the rotation rate is not entirely stable. (The philosopher Immanuel Kant in 1754 was the first to argue that the oceanic tides act as a brake on the Earth's rotation rate; he was unaware that there are other processes taking place on Earth that act to a lesser extent but in the opposite sense.) How to find a better clock? Given any two clocks that do not tick in perfectly strict synchrony, it can be concluded that at least one of them is inaccurate. But which is better?

In his celebrated 1905 paper on special relativity, Einstein seemed to define time in the sense of duration as that read off by ideal clocks. Of course this raises the question as to how we know an ideal clock when we see one or even whether they exist at all. Isaac Newton thought of a free particle moving with uniform speed in a straight line – in accordance with his first law of motion – as a kind of clock. Equal distances are traversed in equal times! (Although non-periodic, it is a clock nonetheless; other examples are radio-carbon dating and Galileo's water clock.) But Newton was fully aware that free inertial motion is an idealisation; no particle is completely free from the influence of the rest of the Universe. As he wrote in the *Principia*: “*It may be, that there is no such thing as an equable motion,*

whereby time might be accurately measured. All motions may be accelerated and retarded, but the flowing of absolute time is not liable to any change.” This is a useful warning. Even the most accurate atomic clocks today suffer from both systematic and unpredictable perturbations [7]. Of course Newton was relying on *theory* to arrive at this pessimistic conclusion: his own theory of universal gravitation. And this reliance on theory is the clue to understanding what a clock is.

FitzGerald himself had already taken a different position from the later one of Einstein. FitzGerald argued that a clock is defined by duration, not the other way round. Of course this raises the question of what duration is and here there is no escape from theory. Put in its simplest terms, a clock is a device which marches in step with the temporal metric picked out in the dynamical equations describing the non-gravitational interactions – electromagnetic, weak and strong nuclear forces. One of the miracles of modern physics – or at least it appears so in the absence of a theory unifying all these interactions – is that they are described by equations that all conform to *the same standard of time*. Modern clocks constructed using either nuclear or atomic frequencies do not appear to be tracking different time metrics even if their accuracies may differ somewhat for technological reasons.

The basic idea of using the stability of atomic frequencies for clocks goes all the way back to Maxwell [6], who realised that one atom is exactly identical to another of the same species, so all devices based on a certain atom’s natural frequencies will in principle be in agreement. But it was not until the 1950s that such technology was developed and since then the accuracy of atomic clocks has improved by roughly a factor of 10 every decade. Today the most accurate are the optical atomic clocks, one of which is estimated to lose a second over 15 billion years, roughly the age of the Universe. As the MIT physicist Daniel Kleppner emphasised in 2006 [8], atomic clocks “*need a lot of love and close personal attention*”. They involve a complex apparatus producing an oscillatory signal that is tuned to be in synchrony with the atoms’ natural oscillations – and counting cycles of this signal is what makes the apparatus a clock. Considerable effort is needed to overcome various kinds of errors that can occur in the process of correcting for perturbations induced by the environment [6].

But the key point is that it is ultimately atomic (quantum) theory that determines the stability of atomic or nuclear frequencies and then there is all the physics involved in designing the oscillator which theoretically locks into and counts these frequencies. Clocks are built to track duration which in turn is defined by our best theories of matter. FitzGerald’s understanding of metric time (later defended by such thinkers as the great French mathematician-physicist Henri Poincaré and more recently Julian Barbour [2], p 170; [9]) was closer to actual practice than perhaps was Einstein’s in 1905.

In his 1902 book *La Science et l’hypothèse*, Poincaré stressed another feature of duration. Suppose one considers two successive periods of time and asks whether they individually correspond to the same duration: “There is no absolute time. When we say that two periods are equal, the statement has no meaning, and can only acquire a meaning by convention. ... Not only have we no direct intuition of the equality of two periods, but we have not even direct intuition of the simultaneity of two events which occur in two different places.” The issue of simultaneity will be taken up in the next section. What did Poincaré mean by “convention” in relation to periods of time passing at a single place in space? His point is that it is not just our best theories of matter that are involved in defining duration, it is finding the mathematical expression of these theories *in their simplest form*. It comes down to choosing temporal coordinates that most simplify the fundamental dynamical equations for the non-gravitational forces. No physicist should go to jail for using arbitrary coordinates, for the end results are the same. But constructing clocks that track the durations defined by arbitrary coordinates would generally be a nightmare. Time is the great simplifier.

4. Spreading time through space: simultaneity

Clocks in daily life normally have two functions. One is to tell the time. We look at the “phase” of a clock – typically the position of its hands – to determine the time of day. The other is the ability to act as a “stopwatch”, thus “measuring” or counting time passing between certain events. So far this paper has been discussing how ideal clocks do this when the relevant events take place where the clock is. Now the question arises as to how the phases of collections of “ideal” clocks separated in space are

correlated. In simple parlance, how does one spread time (clock phase) across space? How does one establish that distant events are simultaneous or not?

If one flies from Auckland, New Zealand, say, to anywhere in the Americas, one arrives before one left, because one has crossed the international date line. The reason for this experience of backwards time travel – which of course has no more deleterious effect on one’s body than a bad case of jet-lag – has to do with the rotation of the Earth. One would not want clocks in New Zealand’s morning to read the same time as clocks in Los Angeles at mid-afternoon. But where one put the international date line is clearly a convention, not rooted in physics. Even where the discrete time zones are delineated (their boundaries often partially deviating from lines of longitude) is largely conventional too (see figure 2). So is there an *objective fact of the matter* whether given two instantaneous events, one occurring say in Rio de Janeiro and the other in London, are simultaneous or not?



Figure 2. Times Zones. Time zones remind us: that we spread time across space on grounds of convenience; that as a result we may end up travelling backwards in time; to ask what the “true” simultaneity is behind the convention.

Sometimes in textbooks on special relativity one sees a picture of a discrete grid in three-dimensional space with ideal synchronised clocks sitting at each grid point, assigning a time to each grid point. Implicitly the continuous limit of this picture in which the grid becomes ever more fine-grained is supposed to represent the nature of spacetime. This is arguably misleading in a number of ways. Firstly, when one looks at a typical experimental laboratory in physics, one does not observe anything like a collection of separated synchronised clocks hanging around. Secondly, and more importantly, it tends to suggest that it is clocks that confer time on the spacetime manifold. This confusion was discussed above. Thirdly, one might think that the problem of defining simultaneity is that of how we synchronise distant clocks and indeed this is a widespread conception. But the real issue *in practice* is how we introduce and justify the temporal coordinates we associate with distant events within spatio-temporal coordinate systems that are best adapted to a given theory. After all, when it is said that two distant events are simultaneous, in physics one means that the “appropriate” temporal coordinates of these events coincide, whether or not there are clocks situated at these events. If there are, the synchrony relations between the clocks are parasitic on the choice of privileged temporal coordinates, not the other way round. And the coordinate systems chosen are themselves parasitic on our best theories.

Consider the case of Newtonian mechanics. Newton’s first law of motion – the law of inertial motion for free particles – does not pick out a privileged simultaneity relation between events. It is only when the second law is introduced along with the inverse square law for the gravitational force that privileged temporal coordinates are singled out. These are such that the gravitational action-at-a-distance appears to be isotropic (the same in all directions) and instantaneous ([8], section 2.2.3.). As Henri Poincaré pointed out in his remarkable 1898 essay *The Measure of Time*, temporal coordinates could be chosen that make the gravitational action retarded in a given direction and backwards-in-time retarded in the reverse direction – with no change in empirical predictions of planetary motion. In that sense the way

we spread time through space in choosing a particular coordinate system is again conventional, but of course, again not all such conventions are equally ‘nice’. No one would choose to work with such awkward non-standard coordinates, least of all Poincaré; they would involve “laws as precise as that of Newton, although more complicated.”

In the case of relativistic field theories, there exists a finite *two-way (there and back)* isotropic invariant speed – the speed of light in the case of electrodynamics. This is confirmed empirically using single clocks. One might naively think that the *one-way* light speed in any direction is measurable in the same sense, therefore objective and not conventional. But determining the one-way speed of any object or signal requires the measurement of the familiar coordinate time – that associated with the time the object takes to traverse a certain distance – and this requires the existence of synchronised clocks at the starting and end points. Again one faces the question: what is the rule for synchronisation? Physicists adopt what has come to be known as the *Poincaré-Einstein convention*, often without thinking. This is the choice of temporal coordinates that make the one-way light speed isotropic or the same in all directions. Again they are not forced to do this: it is a question of formulating the dynamical equations governing the fields in as simple a way as possible. So it is the Poincaré-Einstein convention for spreading time through space that determines whether the distant clocks are synchronised, not the other way round.

There has been a great deal of discussion in the philosophical literature as to whether distant simultaneity is really conventional or not. To some extent this depends on what one means by “conventional”; it surely does not mean arbitrary in the present context! There seems to be less worry that the time (duration) read by single clocks is conventional and this is somewhat surprising, particularly in the light of Poincaré’s insights mentioned in the last section. In his 1989 essay he wrote: “*The simultaneity of two events, or the order of their succession, the equality of two durations, are to be so defined that the enunciation of the natural laws may be as simple as possible.*” Time is the great simplifier, in relation to *both* duration and simultaneity. However, the fact that such simplification can be found in describing the workings of Nature should not be taken for granted. It is decidedly *not* a convention.

5. The relativity of simultaneity

In the process of developing his special theory of relativity, Einstein was drawing out from his postulates the nature of the linear transformations of spatio-temporal coordinates associated with pairs of different, relatively moving, inertial frames of reference. At first he was disturbed by the fact that the transformation of the temporal coordinates means that pairs of events that are simultaneous with respect to one frame are not simultaneous with respect to a relatively moving frame. Eventually under the influence of the Scottish empiricist philosopher David Hume and later the relevant physics of Henri Poincaré, Einstein realised that the conventionality of distant simultaneity *relative to a given frame* means that simultaneity need not be invariant under *change of frames* and this proved to be the key that gave Einstein full confidence in the nature of his transformations as he understood them. (It is often overlooked that Einstein was the first to give the so-called Lorentz transformations – which were first written down by Joseph Larmor in 1900 – the operational significance in terms of the behaviour of rulers and clocks that is taken for granted today; see [10], p 81.)

The relativity of simultaneity in special relativity may not be very intuitive, but it is an undeniable feature of the theory *in its usual guise*. It is true that there are non-standard conventions for spreading time through space that one can adopt in special relativity that make the relativity of simultaneity disappear entirely without changing any of the theory’s empirical predictions. But they are unnatural conventions in the sense that the equations of electrodynamics, for example, become unseemingly complicated – with ugly anisotropies compensating for the introduction of the new conventions. (Such equations will no longer take the same form in all inertial frames, but this is not a violation of the relativity principle because not all frames can use the same simultaneity convention.)

Let us go back to the foundational notion of time outlined in the first section above: a multiplicity of states of affairs of the world divided into moments of time. A question that has exercised a number of

philosophers is whether all such instantaneous states of affairs exist or whether existence must be limited to the “present” state of the world: past and future events did and will exist respectively, but do not now. Perhaps this is what one means by the ‘flow’ of time. In so far as the distinction between the “block Universe” and “presentist” views of time makes sense, it is pretty clear that the relativity of simultaneity creates a difficulty for the latter. The set of events simultaneously taking place in the Universe relative to one inertial observer at a given time will not be simultaneous for a relatively moving observer, even if instantaneously the two observers are in the same place. The option of choosing one observer as privileged seems to violate the relativity principle. Much ink has been spilt in the philosophical literature over this issue, but it is worth bearing in mind that special relativity is not the last word.

In Einstein’s general theory of relativity which he developed after 1905 – the aim being to incorporate the gravitational interaction into relativistic physics – his field equations for gravity do not suggest a relativity of simultaneity in anything like the way it emerges in special relativity. In fact, some theorists have claimed that the dynamical nature of the field equations actually call for a privileged simultaneity relation, linked to what is known as “York time”. (In technical jargon, it is a foliation of spacetime into slices of constant extrinsic curvature [11]). It is true that if one introduces non-gravitational interactions into general relativity, it is standardly assumed that at every point in spacetime there are coordinate systems relative to which the equations for these interactions (mostly) take their special relativistic form. (For explanation of the “mostly”, see [12], section IV.) But such coordinate systems are inertial only in a “local” sense and the relevance of all this to the block Universe versus presentism debate is harder to discern.

6. Time dilation

Suppose there are two ideal synchronised (according to the Poincaré-Einstein convention) clocks at different places on one’s desk (one is an inertial observer) which mark the time of event A taking place at one clock and the later event B at the second clock (using the phase function of the clocks). Let us call the difference between the phase readings the “coordinate time” between the events. Now one imagines a third inertial clock moving in a straight line from location A to location B in just such a way that it can measure the time between these events, using its stopwatch function. Let us call this the “proper time” between the events. According to Einstein’s theory of special relativity, the coordinate and proper times do not strictly coincide – their ratio depends on the square of the ratio of the speed of the moving clock and that of light. In everyday life, the discrepancy is too small to notice. But the ether theorists of the late nineteenth century already had a premonition of such time dilation before Einstein ([10], chapter 4).

It is sometimes said in special relativity that “moving clocks run slow”. But if this means comparing a moving clock against a single stationary clock, care must be taken in evaluating this claim. The standard definition of time dilation involves *three* inertial clocks, comparing a *single* moving clock’s reading against the difference in phases associated with *two* separated, synchronised stationary clocks, as above. And we note that the degree of time dilation (the ratio of proper to coordinate time) depends on the synchrony convention adopted for the stationary clocks.

Where the dependence on the synchrony convention disappears is in the famous “twin paradox” (a sad misnomer because there is nothing paradoxical going on here). The usual version of the twin scenario or rather ‘thought-experiment’ involves a stationary inertial observer who stays at home while their non-inertial twin leaves the Earth on a fast rocket and later returns to the same spot – “non-inertial” because somewhere in the journey the travelling twin has to turn around and therefore decelerate and accelerate back. And even if the vast majority of the trip involves inertial motion, the age of the intrepid travelling twin is according to special relativity less than that of the lazy stay-at-home twin – a consequence of time dilation. The twin effect or “clock retardation” was most famously corroborated in a 1977 experiment involving unstable elementary particles (muons) moving close to the speed of light in a ring accelerator [13]. Today time dilation due to speeds of less than ten metres per second can be detected by atomic clocks.

Given the relativity principle, one might think that from the point of view of the travelling twin, it is the other twin who is moving and hence it is the latter who should be ageing less! But this *reductio ad absurdum* overlooks the fact that there is an asymmetry between the twins; only one suffers accelerations, i.e. a change of inertial rest frames. At least this is the usual argument. But even if both twins travel with identical deceleration and acceleration regimes, differential ageing or clock retardation can occur. For example, suppose both twins leave in identical rockets but one returns before the other and waits for their reunion. Special relativity predicts that the less travelled twin will have aged more when the twins meet. (Not all physicists have been comfortable with clock retardation. An interesting set of letters in *Physics Today* concerning how to understand the twin “paradox” can be found in [14].)

Finally, let us recall that the 1977 experiment mentioned above involves *accelerating* muons. Muons act like clocks because they have a stable half-life. Generally when a clock is taken and accelerated, which means applying a force on it, one should expect trouble. As Sir Arthur Eddington remarked, “*We may force it into its track by continually hitting it, but that may not be good for its time-keeping qualities.*” Let us think of the problem maritime navigators faced for centuries in failing to determine accurate longitudes until John Harrison came along with his famous spring based time-piece, the time-keeping of onboard clocks was corrupted by the heaving motion of the ship. By 1977 there was plenty of evidence that the half-lives of unstable elementary particles were affected to a negligible extent by moving them from laboratory to laboratory, but in the case of the muon experiment the acceleration was 10^{18} times g , the acceleration due to gravity on Earth! The fact that the experiment was consistent with the prediction of time dilation in special relativity demonstrated that *only the instantaneous velocity, not acceleration, contributed sensibly to the retardation effect*. Clocks which meet this constraint satisfy what is often called the “Clock Hypothesis”. In textbooks on special and general relativity it is usual that ideal clocks are implicitly or explicitly assumed to obey the constraint, but such clocks are exceptional. Even atomic clocks will malfunction if not disintegrate when made to accelerate too much. In a sense physicists in 1977 were lucky to find clocks (muons) that withstand accelerations of $10^{18}g$. It took another decade before a calculation was done using perturbation techniques in the theory of the weak interaction to show that the effect of this acceleration on muon half-life was indeed too small to be detectable in the experiment [15]. In general although *uniform velocity produces a universal time dilation, the effect of acceleration depends on the constitution of the particular clock*.

7. Gravitational red-shift

Gravitational time dilation is the last of the ‘properties of time’ that Einstein discovered. In the course of developing his general theory of relativity, he realised that ideal clocks situated at different heights in a gravitational field tick at different rates. In relation to the Earth, the prediction is that as a clock is elevated it goes faster by about one part in 10^{16} for each metre its altitude is increased. For reasons not gone into here, this effect is often called the gravitational red-shift. (The GPS system has to take into account both motion and gravitational time dilation; comparing rates of atomic clocks in different Earth-bound laboratories does too! Gravitational redshift over a difference of height in the Earth’s gravitational field as low as 33 cm has been detected using atomic clocks [7].)

An interesting feature of Einstein’s discovery is that it occurred in 1911, four years before he had found the final form of his gravitational field equations. In fact prediction of the effect, unlike that of the deflection of starlight by the Sun, for example, does not require appeal to the curvature of spacetime. It is enough to appeal to Einstein’s (weak) principle of equivalence: that a homogeneous gravitational field is equivalent to the effects seen by an observer co-moving with a uniformly accelerating frame (though understanding the meaning of the words “gravity” and “equivalent” in this context requires a little thought; see [9], section II) together with the postulates of special relativity. Even today this point is not always recognised in textbooks.

Section 3 started by recalling FitzGerald’s quest to find a definite standard of time and in a sense the issue is returned to when contemplating the significance of gravitational time dilation. As has been seen, atomic clocks have now achieved an astounding degree of accuracy, but no two atomic clocks tick at exactly the same rate and so comparison between such clocks in different parts of the world is called

for. Because of gravitational time dilation, the heights of the laboratories containing the atomic clocks need to be known. This height is not above sea-level but the hypothetical “geoid” which represents a surface of constant gravitational potential. To a first approximation it is a spheroid with the major and minor axes of Earth, and present uncertainty in its location causes an uncertainty of about 3 parts in 10^{17} in relative rates of atomic clocks. But even reducing this uncertainty using more accurate atomic clocks in the future will not resolve the issue. Similar uncertainty is caused by the very dynamic nature of the surface of the Earth, due to solid tidal fluctuations, redistribution of water and other uncontrollable factors. As Daniel Kleppner wrote in 2006, “*The geoid does not lie at rest but is tossed around by a host of processes. ... At the level of accuracy of parts in 10^{17} or 10^{18} , comparing clocks scattered around the world would be no more meaningful than comparing the rates of pendulum clocks on small ships scattered in the oceans, each bobbing its own way and keeping its own time. Which clock could be selected to be the keeper of ‘true’ time? The answer is ‘none’.*” [8] Such clocks are, as Kleppner nicely said, “*too good to be true*”. (He raised the possibility that the primary standard of atomic frequencies could be determined by atomic clocks located in space, but concluded that this “*would not overcome the problem of comparing time or frequency at different locations on Earth.*”)

8. Some reflections

The reader may have noticed that time as an entity in itself has appeared very little in this paper, as opposed to temporal coordinates and the nature and behaviour of physical clocks. Both reflect the fact that physical processes related to the non-gravitational interactions all march to the same tune, *according to our best theories in their simplest form* – the miracle mentioned in Section 3. Clocks whose construction depend on this miracle and in the case of atomic clocks on the sameness of different atoms of a given species are not trying to ‘feel’ some background time but to synchronise with the temporal parameter in our best theories of matter. Perhaps this makes the two phenomena of time dilation a little easier to accommodate conceptually. Motion-induced time dilation ultimately occurs because the equations describing the mechanisms of clocks universally satisfy a symmetry principle – Lorentz covariance. Gravitational time dilation is perhaps harder to fathom at first and Einstein himself was at pains to justify it in 1911, admitting that it seemed “*an absurdity*”. But again it comes down to how accurate clocks comparatively do their job when accelerating at the same rate and separated in the direction of their acceleration. If there is an agent behind all this called time, it is a reclusive, ghost-like beast. As Julian Barbour wrote in 1999, “*It is impossible to understand relativity if one thinks that time passes independently of the world. ... Relativity is not about an abstract concept of time at all: it is about physical devices called clocks.*” ([2], p 137.)

But then there is the important compound notion of *spacetime* in relativistic physics and in particular in general relativity wherein certain gravitational effects are often attributed to the ‘bending’ or curvature of spacetime due to the presence of mass. It should be borne in mind, however, that Einstein himself did not consider his theory to be geometrical in any fundamental sense [16]. His mature view (after considerable soul-searching) was that it was a physical field – the “metric” field – and its contortions in the presence of mass that are responsible for these effects. And as he famously said, “*Fields are not defined on space, space is a structural property of fields*” – and this goes for spacetime too. Perhaps Newton would have approved. (For wide-ranging discussions of time in physics and the difference between space and time, see [2], [3] [17] and [18].)

Acknowledgements

Much of my thinking about time has been influenced by Julian Barbour’s extensive writings on the subject and by lively discussions with him over many years. I thank Dennis Lehmkuhl and James Read for discussions about time in general relativity. Finally, thanks go to Jo Ashbourn for creating the HAPP conference lectures and inviting me to contribute to them.

References

- [1] Feynman R P, Leighton R B and Sands M 1965 *The Feynman Lectures on Physics, Volume 1*

- (Reading, Mass.: Addison-Wesley)
- [2] Barbour J 1999 *The End of Time. The Next Revolution in Our Understanding of the Universe* (London: Weidenfeld and Nicolson)
 - [3] Barbour J 2020 *The Janus Point. A New Theory of Time* (London: The Bodley Head)
 - [4] Mackey M C 2003 *Time's Arrow: The Origins of Thermodynamic Behavior* (Mineola, New York: Dover Publications Inc.)
 - [5] Tegmark M 2014 *Our Mathematical Universe: My Quest for the Ultimate Nature of Reality* (Alfred A. Knopf Inc.)
 - [6] Barbour J, Koslowski T and Mercati F 2014 *Phys. Rev. Lett.* **113** 181101
 - [7] Ludlow A D, Boyd M M, Ye J, Peik E and Schmidt P O 2015 Optical atomic clocks *Rev. Mod. Phys.* **87** 637 (*Preprint* atom-ph/1407.3493v2)
 - [8] Kleppner D 2006 *Phys. Today* **59**(3) 10
 - [9] Barbour J 2009 The nature of time *Preprint* gr-qc/0903.3489v1
 - [10] Brown H R 2007 *Physical Relativity. Spacetime Structure from a Dynamical Perspective* (Oxford: Oxford University Press)
 - [11] Roser P 2016 *PhD thesis* Clemson University *Preprint* gr-qc/1609.03942v1
 - [12] Brown H R and Read J 2016 *Am. J. Phys.* **84**(5) 327
 - [13] Bailey J *et al.* 1977 *Nature* **268** 301
 - [14] Terrell J *et al.* 1972 *Phys. Today* **25**(1) 9
 - [15] Eisele A M 1987 *Helv. Phys. Acta* **60** 1024
 - [16] Lehmkuhl D 2014 *Stud. Hist. Phil. Mod. Phys.* **46** 316
 - [17] Callender C 2017 *What Makes Time Special?* (Oxford: Oxford University Press)
 - [18] Lockwood M 2005 *The Labyrinth of Time. Introducing the Universe* (Oxford: Oxford University Press)