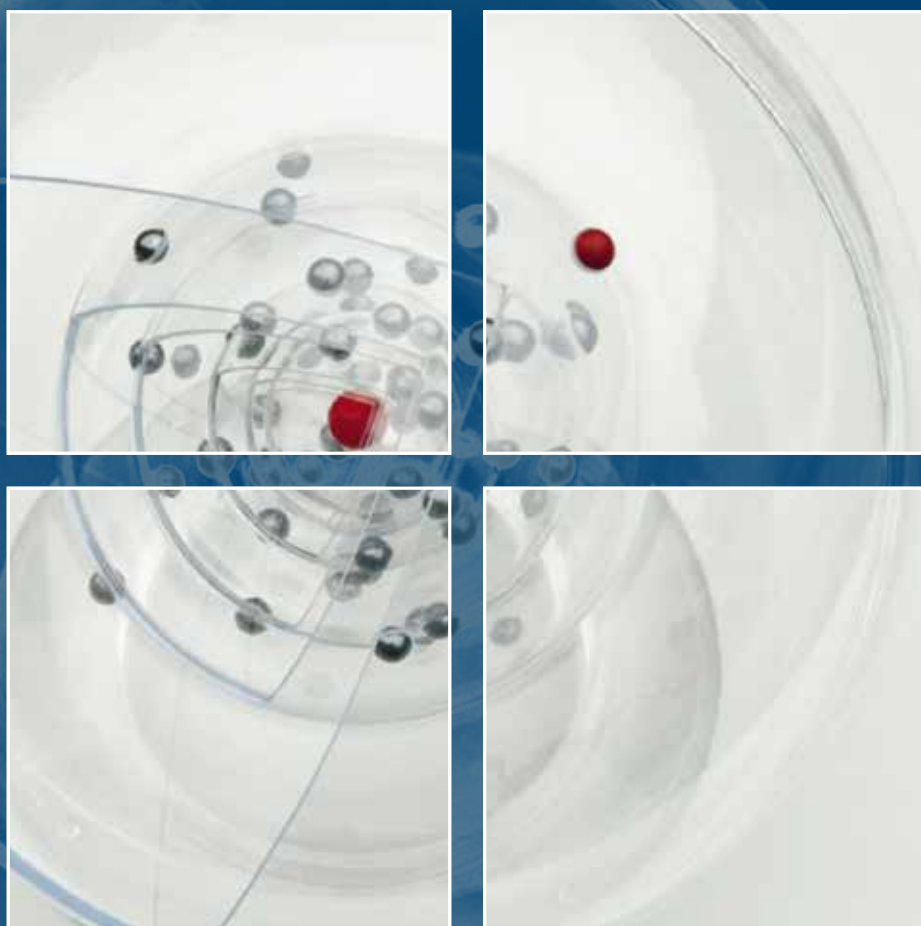


50 Jahre atomare Definition der Sekunde

50 Years of the Atomic Definition of the Second



**Fachorgan für Wirtschaft und Wissenschaft, Amts- und
Mitteilungsblatt der Physikalisch-Technischen Bundesanstalt
Braunschweig und Berlin**

127. Jahrgang, Heft 3, September 2017

50 Jahre atomare Definition der Sekunde

50 Years of the Atomic Definition of the Second

Titelbild:

Plexiglasmodell eines Caesiumatoms. Neben dem Atomkern ist auch das äußerste Elektron, das für den Uhrenübergang verantwortlich ist, farbig markiert. Dieses Modell wurde im Jahr 2005 anlässlich einer Ausstellung im Braunschweigischen Landesmuseum im Auftrag der PTB gebaut. Titel der damaligen Ausstellung: „Rezept für eine Atomuhr“. Ein paar Jahre später machte dieses Modell auch Werbung für die Ausstellung „Caesium“ des Künstlers Evarist Richer im Haus Salve Hospes des Kunstvereins Braunschweig. Quelle: PTB.

Inhalt – Contents

50 Jahre atomare Definition der Sekunde – 50 Years of the Atomic Definition of the Second

- Einführung 03
- Die zweite Teilung der Stunde. Zur Geschichte der Sekunde 05
Johannes Graf
- Atomare Definition der Zeiteinheit 1967–2017 13
Andreas Bauch
- NPL's Contribution to the Introduction of the Caesium Second 33
Peter Whibberley
- A Historical Review of U. S. Contributions to the Atomic Definition of the SI Second 41
Michael A. Lombardi
- Atomuhren und Navigation mit Satelliten-Systemen (GNSS) 53
Gerhard Beutler
- Einstein und die Zeit 65
Claus Kiefer
- Auf dem Weg zu einer Neudefinition der Sekunde: Was kommt nach Caesium? 73
Ekkehard Peik

PTB-Innovationen

- Ausgesuchte Technologieangebote 79

Einführung

Jörn Stenger*

Die Einheit der Zeit ist neben den Einheiten für Länge und Gewicht die historisch wohl wichtigste Einheit – unmittelbar erfahrbar und für das tägliche Leben notwendig. Es ist mehr als lohnenswert, sich anlässlich des fünfzigsten Jahrestages der aktuellen Definition der Einheit Sekunde mit der Historie, den Auswirkungen der Definition und möglichen zukünftigen Entwicklungen zu beschäftigen. Momentan steht das gesamte Einheitensystem (SI) vor einem Umbruch, und das hat viel mit der Definition der Sekunde zu tun.

Diese Definition ist, wie die für alle Basiseinheiten, per Beschluss eines internationalen Gremiums, der Generalkonferenz der Meterkonvention, festgelegt worden. So ist jede Veränderung im physikalischen Einheitensystem auch eine politische Angelegenheit – mit weitreichenden Auswirkungen. Als die Generalkonferenz im Jahr 1967 beschloss, die Sekunde auf der Grundlage eines atomaren Quantenüberganges zu definieren, machte es diese Einheit zum Vorreiter einer zukunftsweisenden Entwicklung, die eine Einheit nach der anderen erfasste und im kommenden Jahr vermutlich das gesamte Internationale Einheitensystem revolutionieren wird: Dann wird nämlich jede Einheit auf der Basis eines festgelegten Zahlenwertes einer Konstanten ruhen.

Die Zeit an sich ist von der Natur vorgegeben, sie ist in vielen physikalischen Gesetzmäßigkeiten mit anderen Größen verknüpft. Die Zeiteinheit Sekunde jedoch ist eine Übereinkunft zur Messung bzw. technischen Darstellung der Zeit. Damit besteht die Freiheit, diese Einheit geeignet festzulegen. Eine Einheitendefinition legt lediglich die numerischen Faktoren fest, mit denen wir Messwerte beschreiben, und man wählt sie sinnvollerweise so, dass typischerweise vorkommende Messwerte mit handhabbaren Zahlen beschrieben werden können. Die Festlegung von 1967 war von

daher willkürlich, man hat sie aber praktischerweise so gewählt, dass alte und neue Sekunde möglichst gut übereinstimmten. Warum hat man nun damals diesen etwas unanschaulichen Quantenübergang gewählt?

Jedes Einheitensystem richtet sich am Bedarf der Zeit aus. Was die Zeit angeht, so war es über Jahrhunderte wichtiger, sie überall einfach bestimmen zu können – wichtiger, als sie mit höchster Präzision zu kennen. Was ist da nützlicher, als die Bewegung der Erde gegenüber dem Fixsternhimmel zu verwenden? Doch für moderne Anforderungen reichte dies nicht mehr aus. Bereits Maxwell und vor allem Max Planck haben vor mehr als hundert Jahren darauf hingewiesen, dass universelle Naturkonstanten die idealen Referenzen für Einheitendefinitionen darstellen. Man braucht dann lediglich noch eine Apparatur, die eine Konstante in einer ungestörten Form messbar macht, und erhält eine universelle, stabile und im Prinzip beliebig genaue Realisierung der Einheit der Konstante. Mit der Neudefinition der Einheit Sekunde kam man diesem Ziel sehr nahe: Der Caesiumübergang stellt ein relativ einfaches Quantensystem dar, das bereits damals experimentell gut zugänglich war. Heute sind mögliche Störungen des Systems sehr gut verstanden und lassen sich beherrschen, so dass man Messunsicherheiten im Bereich von 10^{-16} erreicht. Das ist weitaus genauer als bei jeder anderen Einheit.

Die Einheitendefinition der Sekunde war eine Initialzündung für andere Definitionen: So wurde 1983 die Einheit Meter an die Sekunde gekoppelt, indem man den Zahlenwert der Lichtgeschwindigkeit festlegte und so in Kombination mit dem Caesiumübergang die Einheit Meter definierte. Das war ganz im Sinne der Planck'schen Forderung, denn die Lichtgeschwindigkeit ist eine fundamentale, unbeeinflussbare Naturkonstante. Dieser Weg

* Dr. Jörn Stenger, Abteilung „Ionisierende Strahlung“, Mitglied des Präsidiums, E-Mail: joern.stenger@ptb.de

wurde mit der Einführung der quantenelektrischen Einheiten unter Bezug auf Festlegungen der Zahlenwert für die Planckkonstante und Elementarladung weitergegangen, momentan allerdings noch formal außerhalb des gültigen SI-Systems. Das wird sich ändern, wenn die Generalkonferenz der Meterkonvention wie erwartet Ende 2018 das gesamte SI auf die Basis von Konstanten stellt.

Die Definition der Einheit Sekunde hat somit nicht nur technologische Durchbrüche ermöglicht, wie z. B. die Entwicklung von Satelliten-Navigationssystemen, sondern sich auch vielfältig auf das gesamte Einheitensystem ausgewirkt – ein guter Grund für einen Band der PTB-Mitteilungen anlässlich des fünfzigsten Jahrestages der „Caesium-Sekunde“.

Die zweite Teilung der Stunde. Zur Geschichte der Sekunde

Johannes Graf*

Der berühmte braunschweiger Mathematiker und Naturwissenschaftler Carl Friedrich Gauß hat sich große Verdienste um die Standardisierung von Maßen und Gewichten erworben, denn 1831/1832 entwickelte er zusammen mit Wilhelm Eduard Weber das Centimeter-Gramm-Sekunde-System. Er gab damit die Blaupause für spätere Systeme von Basiseinheiten vor. Schon bei Gauß bildete die Sekunde die Grundlage der Zeitmessung. Doch anders als vor genau fünfzig Jahren, als die SI-Sekunde über die Schwingungen des Caesium-Atoms definiert wurde, hatte Gauß die Sekunde nicht direkt bestimmt, sondern von größeren Zeiteinheiten lediglich abgeleitet. Bis in die 1950er-Jahre galt, was er bereits Anfang der 1830er-Jahre vorgeschlagen hatte: Die Sekunde sei der 86.400ste Teil eines mittleren Tages, also der Drehung der Erde um sich selbst. Wie sich die Definition und Bestimmung der Sekunde im Lauf der Geschichte entwickelte, soll im Folgenden beschrieben werden.

Die Sekunde – ein Spätling

In der Zeitmessung tauchte die Sekunde lange Zeit überhaupt nicht auf. Während der klassischen Antike hatte man sich vorwiegend mit Fragen des Kalenders und der Teilung des Tages in Stunden und vielleicht noch Minuten beschäftigt. Dabei spielten die Bruchteile der Minute praktisch keine Rolle. Auch war die Einteilung der Zeit noch keine feste Norm, sondern wurde in den unterschiedlichen Kulturen ebenso unterschiedlich vollzogen. So speist sich unsere heutige Einteilung des Tages aus diversen Quellen: Die Einteilung der Stunde in 60 Minuten verdankt sich dem babylonischen Sexagesimalsystem, die Zweiteilung des Tages in je 12 Licht- und Nachtstunden dem antiken Ägypten und die gleichlangen Stunden der griechischen Astronomie [1].

Wieso Zeit- und Kalenderangaben bis heute Vielfache von 12 oder 60 sind, dafür gibt es verschiedene Erklärungen. Zum einen spielte sicherlich eine Rolle, dass das Jahr 12 komplette

Mondzyklen umfasst. Die Zahl 12 fand sich so schon in den großen Vorgängen der Natur und des Kosmos wieder. Sie galt deshalb im antiken Babylon als heilige Zahl. Ebenso wichtig war sicherlich auch, dass die Zahlen 12, 24 und 60 hochzusammengesetzte Zahlen sind. Diese sind durch sehr viele Zahlen ohne Rest teilbar. Damit kann man in vielen Fällen die im Alltag eher aufwendigen Bruchrechnungen vermeiden. Auch für Kreisteilungen sind 12, 24 und 60 sehr geeignete Zahlen. Dies mag auch als Erklärung dienen, wieso später die Minute ebenfalls in 60 Sekunden geteilt wurde.

Die mittelalterliche Astronomie als Geburtshelfer

Im frühen Mittelalter nahm die Diskussion um die Bruchteile der Stunde oder gar Minute im Kontext der populären Computus-Texte an Fahrt auf. Als Computus bezeichnet man ein Lehr- und Tabellenwerk, das Möglichkeiten zur Ermittlung des beweglichen Osterdatums beschreibt. Außerdem behandeln viele dieser Texte weitergehende Fragen zur christlichen Zeitrechnung und ihrer praktischen Umsetzung. Sie vermitteln deshalb auch Grundkenntnisse in Astronomie und Mathematik, soweit sie für die Zeitbestimmung notwendig sind.

In den Jahren 703 und 725 verfasste der englische Mönch Bede Venerabilis seine beiden viel gelesenen Texte *De temporibus* und *De temporum ratione*, die wichtig für die weitere Diskussion um die Teilung der Stunde wurden [2]. Intensiv wurde bei Bede und anderen Zeitgenossen die Frage diskutiert, ob es eine kleinste Einheit der Zeit gäbe, oder ob die Zeit ein Kontinuum darstelle, bei dem jede kleine und kleinste Zeitspanne noch beliebig oft geteilt werden könne. Dabei orientierte man sich nicht am Messbaren, sondern am Wahrnehmbaren. Für Bede stellte der Wimpernschlag die kleinste erkennbare Zeiteinheit dar. Er nennt dafür aber keine Vergleichsgröße. Denn die Diskussion über eine „Zeit aus

* Dr. Johannes Graf, stellvertretender Museumsleiter – Wissenschaftlicher Mitarbeiter Deutsches Uhrenmuseum – Hochschule Furtwangen

Atomen“ hielt er für müßig. Für ihn waren alle Teilungen der Zeit von Menschen gemacht. Spekulationen wie die von Hrabanus Maurus in *De computo* (820), der die Stunde in 22.560 Zeitatome ($\sim 0,16$ Sekunden) einteilte, lehnte er ab. Aber noch im Spätmittelalter wurde die Vorstellung einer kleinsten, nicht mehr teilbaren Zeiteinheit ernsthaft diskutiert.

Spätestens im 14. Jahrhundert war in gelehrten Kreisen die 60er-Teilung der Stunde und Minute als Rechengröße weithin bekannt, vor allem durch den Einfluss der islamischen und maurischen Astronomie. Größten Einfluss dabei übte ein astronomisches Tabellenwerk aus, die sogenannten *Alphonsinischen Tafeln*. Verfasst wurden sie zwischen 1252 und 1270 von arabischen, jüdischen und christlichen Gelehrten, die König Alphons X, genannt der Weise, von Kastilien und Leon an seinem Hof in Toledo versammelt hatte. Die astronomischen Tabellen verwendeten nun durchgehend die aus der babylonischen Tradition stammenden Sexagesimalteilungen, die sich besonders gut für Kreisteilungen eigneten. Bogengrade, Tage und Stunden konnten bei Berechnungen in 60 Minuten, diese wiederum in 60 Sekunden zu jeweils 60 Terten geteilt werden. Doch spielten in diesem Tabellenbuch die zweite oder gar dritte Teilung der Stunde für astronomische Angaben

noch keine praktische Rolle. Wichtigste Messgröße für die Beobachtung des Himmels stellten die „*minuta diei*“ dar. Sie entsprachen – wie eine Umrechnungstafel in den *Alphonsinischen Tafeln* mitteilt – zwei Minuten und dreißig Sekunden. Eine Zeitstunde umfasste folglich 24 „*minutae diei*“.

Dohrn-van Rossum hat in seiner „Geschichte der Stunde“ zusammenfassend für das Mittelalter festgehalten, dass bereits Mitte des 14. Jahrhunderts „die Teilung der Stunde in 60 Minuten und der Minute in 60 Sekunden als abstrakter Rahmen für Denken und Handeln üblich geworden sei. Mindestens bis zum Ende des 16. Jahrhunderts kann jedoch vom gewöhnlichen Gebrauch von Minuten oder gar Sekunden keine Rede sein. Minuten- und Sekundenangaben finden sich nur [...] in theoretischen Erörterungen und bei astronomischen oder astrologischen Zeitangaben. Sie waren uralte theoretische, nicht aber messbare Zeiteinheiten“[3].

Das Pendel bringt den Durchbruch für sekundenzeigende Uhren

Die erste Nachricht über eine Uhr mit Sekundenzeiger stammt aus dem Jahr 1557 [4]. Doch zu dieser Zeit konnte man noch keine Zeitmesser herstellen, die auch nur annähernd so präzise waren, dass ein dritter Zeiger an der Uhr sinnvoll war. Erst ein Vierteljahrhundert später gelang es dem genialen Uhrmacher Jost Bürgi, einige Uhren zu bauen, die nicht bloß Minutenbruchteile angaben, sondern eine Zuverlässigkeit und Ganggenauigkeit erreichten, die eine zusätzlichen Sekundenanzeige rechtfertigten [5]. Bürgi hatte die Uhrentechnik an drei entscheidenden Punkten vorangebracht, indem er störende Einflüsse auf den stetigen Gang der Uhr ausschaltete: Zum einen arbeitete er sorgfältiger und genauer als andere zeitgenössische Uhrmacher – insbesondere beim Räderwerk, zum anderen ersetzte er die bislang übliche Spindelhemmung mit Waag durch die bessere Kreuzschlaghemmung. Die dritte Verbesserung betraf die Energiezufuhr der Uhr. Damalige Uhren hatten das Problem, dass die Antriebskraft stark variierte, je nachdem, ob die Feder gespannt oder eher entspannt war. Bürgi gelang es, die Federkraft zum Antrieb weitgehend gleichzuhalten, indem er das Remontoir einführte. Bei diesem Regelkreislauf löste das Uhrwerk in periodischen Abständen ein zweites Federwerk aus, das lediglich dazu diente, der eigentlichen Antriebsfeder Energie zuzuführen. Für den Gleichlauf der Uhr war es unerheblich, wenn im zweiten Werk die Federspannung variierte. Wichtig war, dass durch das regelmäßige Aufziehen nur ein kleiner Bereich der Federspannung zum Betrieb des eigentlichen Uhrwerks genutzt wurde. Dadurch blieb die Antriebsenergie weitgehend konstant.

Bild. 1:
Nachbau des
Astrariums, einer
astronomischen
Uhr mit Darstellung
der Planetenbewegungen von Giovanni Dondi, 1364
[Musée International D'Horlogerie,
La-Chaux-de-Fonds]



Nur wenige Uhren mit Kreuzschlaghemmung wurden gebaut (Bild 2). Das lag sicherlich mit daran, dass man diese Uhren mit Sekunden-Tick nur für ganz spezielle Zwecke benötigte, etwa die Astronomie. Und dabei spielte es kaum eine Rolle, ob diese Uhren auf wenige Sekunden am Tag genau waren. Für die meist kurzen Zeitspannen – etwa im Bereich weniger Minuten – tickten diese frühen Uhren hinreichend konstant.

Erst mit dem Entstehen der klassischen Naturwissenschaften, ab der zweiten Hälfte des 17. Jahrhunderts, benötigte man deutlich mehr Uhren mit Sekundenanzeige, zum Beispiel für physikalische Experimente. Diese Uhren waren aber keine Uhren mit Kreuzschlaghemmung mehr. Weniger heikel in der Anwendung waren Uhren mit Pendel, das Galileo Galilei und Christiaan Huygens entdeckt hatten. Egal, wen man von beiden als eigentlichen Erfinder des Pendels als Regulator der Uhr ansieht, in einem sind sich alle einig: Die Einführung des Pendels im Uhrenbau markiert eine Epochenschwelle. Bis zum Ende des 16. Jahrhunderts hatten die Handwerker allein bestimmt, welche Prinzipien und welche Techniken in den Uhrwerken verwendet wurden. Im 17. Jahrhundert begann die Theoretisierung der Mechanik zur Wissenschaft, zur *nuova scienza* (Galilei). Diesen Prozess der „Verwissenschaftlichung im Handwerk“ prägten nicht mehr die Praktiker, sondern die an den Höfen oder für Akademien arbeitenden Gelehrten [6]. Auch wenn das eigentliche Handwerk des Uhrenbaus die Domäne der Mechaniker blieb, so übten die Wissenschaftler, insbesondere die Astronomen, einen nicht zu unterschätzenden Einfluss auf die weitere Entwicklung der Uhrentechnik aus, sei es mit neuartigen wissenschaftlichen Erkenntnissen oder als Auftraggeber für Präzisionsinstrumente.

Mit der Erfindung des Pendels setzte eine rasante Entwicklung bei der Genauigkeit ein, die innerhalb von wenigen Jahrzehnten die Gangabweichungen radikal reduzierten. Mitte des 18. Jahrhunderts war bei Standuhren eine Abweichung von wenigen Sekunden am Tag üblich. Die Uhr war jetzt nicht mehr nur für eine kurze Zeitspanne genau, wie man sie für wissenschaftliche Experimente benötigte, sondern lief auch auf Dauer weitgehend konstant. Doch dazu hatte es noch einiger grundlegender Entdeckungen bedurft.

Die erste entscheidende Verbesserung beruhte auf der Kombination des Pendels mit der Ankerhemmung in den 1670er-Jahren, eine Erfindung, die dem Wissenschaftler Robert Hooke oder den Uhrmachern William Clement bzw. Joseph Knibb zugeschrieben wird [7]. 1715 gelang es dem Uhrmacher George Graham, das Schwingsystem durch die sogenannte ruhende Ankerhemmung noch einmal grundlegend zu verbessern.

Auch das Pendel selbst wurde mehrfach weiterentwickelt: Das von Huygens formulierte Prinzip des Isochronismus besagt, dass die Frequenz eines Pendels abhängig ist von der wirksamen Länge. Besonders in gemäßigten Breiten hatte man folglich mit einem gravierenden Problem zu kämpfen, wenn im Wechsel der Jahreszeiten die Außentemperaturen stark schwankten. Denn dadurch änderte sich auch die Länge des Pendels und letztendlich der Gang der Uhr. Es war deshalb entscheidend, diesen Einfluss der Temperatur auf das Pendel zu verringern. Unter denjenigen, die auf diese Weise die Genauigkeit der Uhren verbessern wollten, war John Harrison, dessen Name eng mit der Erfindung des Marine-Chronometers und der Navigation auf See verbunden ist. Um 1725 entwickelte Harrison als erster ein sogenanntes Rostpendel. Diese Konstruktion kombinierte Stangen zweier Metalle mit gegenläufigem Ausdehnungskoeffizienten so, dass sich deren Einflüsse auf die Länge des Pendels neutralisierten. Etwa gleichzeitig erfand George Graham eine andere wirksame Konstruktion zur Verminderung der Längenänderung des Pendels. Er nutzte den großen Einfluss der Temperatur auf die Ausdehnung von Quecksilber und dessen hohe Dichte für sein Kompensationspendel.

Bild 2:
„Wiener Kristalluhr“
mit Sekunden-
zeiger, Jost Bürgi,
Prag 1622/1627
[Kunsthistorisches
Museum Wien,
Kunstammer Inv.
KK 1116]



Vor allem in England setzten sich die neuartigen Uhren schnell in bürgerlichen und adligen Kreisen durch. Dabei zeigte sich, dass die traditionelle Zeit der Sonnenuhren mit der von den Pendeluhrn gemessenen nicht leicht in Übereinstimmung zu bringen war. Anfang des 18. Jahrhunderts prallten zwei Zeitkulturen aufeinander.

Traditionell hatte man die öffentlichen Turmuhrn nach Sonnenuhren gestellt, die sich bis heute an den Außenwänden einiger Kirchen erhalten haben. Durch die starken Gangschwankungen der Uhren fiel im Alltag nicht auf, dass die Zeit der Sonnenuhren im Jahresverlauf alles andere als gleichförmig ist. Denn der ellipsenförmige Lauf der Erde um die Sonne und die Schräglage der Achse bedingen, dass die Tageslänge im Jahresverlauf um bis zu zwanzig Sekunden täglich schwankt, je nachdem, an welcher Position der Umlaufbahn sich die Erde gerade befindet. Die täglichen Änderungen summierten sich zu jahreszeitlichen Abweichungen des Sonnenuhrzeigers von bis zu 15 Minuten plus/minus eines Tages mit genau 24 Stunden. Diese Unterschiede entwickelten sich zu einem echten Problem, als die Standuhren konstant jeden Tag 24 Stunden mit je 60 Minuten und 60 Sekunden zählten. Welche Zeit sollte nun gelten, die der Sonnenuhren oder die der Pendeluhrn?

Bild 3:
Standuhr mit Angabe der Wahren Zeit:
„*Horae Indicantur
apparentes Involutis
Aequationibus*“,
William Scafe,
London um 1730
[Deutsches Uhren-
museum Furtwan-
gen, Inv. 2009-054]



Viele Zeitgenossen am Beginn des 18. Jahrhunderts waren es zunächst noch gewohnt, sich nach der Sonnenuhr zu richten. Sie nahmen die Pendeluhrn als Störung ihrer Routine war. Für sie zeigten nicht die Sonnenuhren, sondern die Pendeluhrn die „falsche“ Zeit an. Deshalb baute man vereinzelt Standuhren, die die Unregelmäßigkeiten der Sonnenuhren nachahmten, indem diese die Pendellänge im Laufe des Jahres veränderten (siehe Bild 3). Die Folge daraus: Zwar zerteilten diese Pendeluhrn den Tag mit 24 Stunden zu 60 Minuten und 60 Sekunden in 86.400 Teile – doch waren die Sekunden unterschiedlich lang, je nach Jahreszeit!

Für wissenschaftliche Zwecke waren solche Uhren nicht zu gebrauchen. Um genaue Zeitmessungen vorzunehmen, baute man schon im 18. Jahrhundert Pendeluhrn, welche die gleichmäßige, so genannte „Mittlere Zeit“ anzeigten, daneben aber auch die „Wahre Zeit“, welche auf den Sonnenuhren abgelesen wurde.

Es sollte jedoch noch dauern, bis sich die mittlere Zeit im Alltag durchsetzen konnte. Zwar waren ab dem Ende des 18. Jahrhunderts Tabellenwerke weit verbreitet, die die Umrechnung der „Wahren Ortszeit“ in eine „Mittlere Ortszeit“ auch für Nicht-Astronomen leicht möglich machten, und ihre Zahl nahm in der ersten Hälfte des 19. Jahrhunderts noch stark zu. Doch wie stark die Beharrungskräfte selbst in ausgesprochenen Uhrenmetropolen waren, zeigt das Beispiel von Genf. Erst 1821 hatte man dort beschlossen, an den öffentlichen Uhren die mittlere statt der wahren Ortszeit anzuzeigen [8].

Im 19. Jahrhundert war man sich unter den Gelehrten einig, dass die Sekunde als eine der Basiseinheiten zu gelten habe, obwohl man sie nicht wie den Meter oder das Kilogramm materiell zu fassen bekam. 1831 tauchte die Sekunde auch im Titel des von Carl Friedrich Gauß erarbeiteten Centimeter-Gramm-Sekunde-Systems auf, das sich in der zweiten Hälfte des 19. Jahrhunderts zum Meter-Kilogramm-Sekunde-System wandelte. Während die Metrologen die möglichst exakte Bestimmung des Urmeters und Urkilogramms mit großem Aufwand betrieben, war es merkwürdig ruhig um die dritte Basiseinheit. Bis zur Mitte des 20. Jahrhundert ging man stillschweigend davon aus, dass die Sekunde der 86.400ste Teil eines gleichlangen mittleren Sonnentages war.

Die Bruchteile der Sekunde

Neben der selbstverständlichen Annahme, dass die Sekunde immer gleich lang sei, wenn man sie anhand der Erddrehung bestimmte, war noch etwas Zweites erstaunlich: Die Sekunde fand Eingang in das System von Basiseinheiten, obwohl die Teilung der Zeit ja gerade nicht wie bei Länge

und Gewicht auf dem Dezimalsystem beruhte, sondern auf dem Duodezimalsystem (die Sekunde als der 3600ste Teil einer Stunde).

Anfangs galt das auch für die Bruchteile der Sekunde. Deutlich wird das an Tertienzählern, die seit dem ausgehenden 18. Jahrhundert zur Standardausrüstung von Sternwarten gehörten. Der Name dieser Stoppuhren rührt daher, dass sie nicht bloß die erste Teilung der Minute und die zweite der Sekunde anzeigten, sondern ebenso die dritte in 60tel-Sekunden. Auf dem dritten Zifferblatt werden diese Bruchteile der Sekunde meist jedoch nur mit den Zahlen von 1 bis 6 bezeichnet, denn die verwendete Uhrentechnik begrenzte die Taktfrequenz des Schwingensystems.

Mitte des 19. Jahrhunderts entstanden erstmals Messgeräte, mit denen noch kleinere Zeitabschnitte gestoppt werden konnten. Diese Kurzzeitmesser wurden von ihrem Erfinder Matthäus Hipp als Chronoskope bezeichnet (Bild 4). Die neuartigen Stoppuhren zergliederten die Sekunde in Zehntel, Hundertstel oder Tausendstel. Diese dezimale Zählung der Bruchteile von Sekunden setzte sich bald anstelle der 60er-Teilung durch.

Im Gefolge der Französischen Revolution, in den 1790er-Jahren in Frankreich, stellte man die Zeit für einige Jahre auf das Dezimalsystem um [9]. Doch konnte sich diese neue Einteilung nicht auf Dauer halten. Ähnliche Diskussionen im Umkreis der Meterkonvention in den letzten Jahrzehnten des 19. Jahrhunderts fanden dann nur noch auf dem Papier statt. Nur für wenige Anwendungen, wie den Studien der Arbeitsabläufe in der Industrie, bürgerte sich ein, Uhren zu verwenden, die statt der Sekunde 100stel-Minuten zeigten (Bild 5). Im Alltag jedoch hält man bis heute bei der Tageseinteilung in 24 Stunden zu 60 Minuten und 60 Sekunden fest.

Quarzuhren holten die Zeit vom Himmel

Es deutete sich bereits mit Pendeluhr an, die von William Shortt konstruiert und bei der britischen Firma Synchronome in den 1920er-Jahren gebaut wurden: Etwas konnte mit der Erddrehung nicht stimmen.

Die Shortt-Synchronome-Uhren erreichten eine bis dato unschlagbare Präzision durch ihre raffinierte Kombination aus zwei Pendeluhrn, bei denen die eine die andere auf elektrischem Wege steuerte [10]. Sie waren bereits auf einige zehn Millisekunden am Tag genau. Doch immer wieder schien die gemessene Zeit der Shortt-Uhren von der astronomisch bestimmten abzuweichen, selbst wenn man die bekannten Unregelmäßigkeiten berücksichtigte – sei es die kontinuierliche Abbremsung der Erde durch die Gezeiten gegenüber dem Mond und die daraus resultierende Verlängerung des Tages (und auch der Sekunde) oder



Bild 4:
Chronoskop mit Angabe der 1/1000-Sekunde, Matthäus Hipp, Neuchâtel um 1867 [Deutsches Uhrenmuseum Furtwangen, Inv. 2015-068]



Bild 5:
Stoppuhr mit Dezimalskala für 1/100-Minuten, ISGUS, Schwenningen, um 1920 [Deutsches Uhrenmuseum Furtwangen, K-0747]

das Trudeln der Erde um ihre Achse. Die Astronomen vermuteten, dass es zusätzliche Einflüsse gab; nachweisen konnte man sie jedoch noch nicht. Dies sollte erst mit Quarzuhren gelingen.

Die ersten elektronischen Uhren waren Nebenprodukte der boomenden Hochfrequenztechnik. Quarzuhren wurden von den gleichen Leuten erfunden, die für Industrie, Militär und Wissenschaft Frequenznormale und Messgeräte entwickelten. Da die Frequenz bekanntermaßen das Inverse der Zeit ist, konnte man diese Geräte eichen, indem man sie mit der Uhrzeit verglich. Deshalb lag es im Umkehrschluss nahe, Hochfrequenznormale auch bei Uhren zu verwenden.

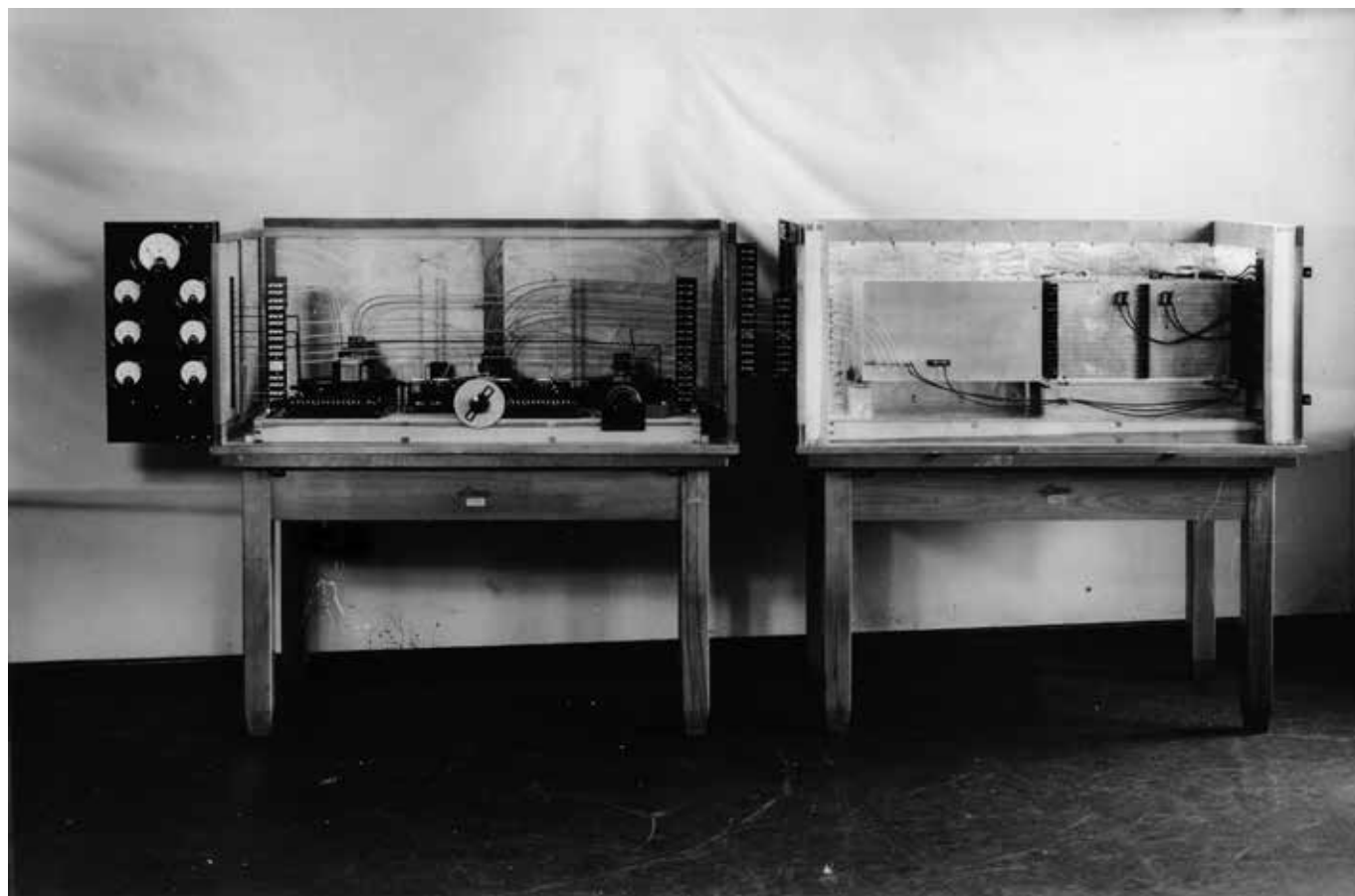
1928 baute Williamarrison bei den US-amerikanischen Bell Labs die erste elektronische Uhr mit Schwingquarz, die im Bereich von Kurzzeitmessungen selbst den allerbesten bekannten Zeitmessern haushoch überlegen, auf Dauer jedoch den Präzisionspendeluhren unterlegen war. Kein Wunder: Denn die Erbauer der neuartigen Uhren interessierte es nicht, wie spät es war, sondern wie lange etwas dauerte. Dieser Blickwechsel sollte die Zeitmessung ähnlich tief erschüttern wie die Relativitäts- und Quantentheorie die Physik. Denn bereits die ersten Quarzuhren leiteten die Sekunde nicht mehr wie die Astronomen durch Teilung des Tages her, sondern durch Frequenzteiler, die eigentlich elektronische Schaltkreise zur Vervielfachung kleinster Zeitspannen darstellen.

Fieberhaft arbeitete man in den nächsten Jahren daran, diese neuartigen Uhren auch für die Zeitmessung über einen längeren Zeitraum nutzen zu können. 1932 gelang es Adolf Scheibe und Udo Adelsberger in der Physikalisch-Technischen Reichsanstalt Berlin, eine auf Dauer zuverlässige und praktikable Technik zu entwickeln [11]. Ab diesem Jahr liefen vier elektronische Zeitmesser langzeitstabil mit einer Genauigkeit, die auch die Shortt-Uhren noch um ein Vielfaches übertrafen.

1934 stellten die Ingenieure in Berlin fest, dass alle vier Quarzuhren auf exakt gleiche Weise um einen Hauch von der amtlichen, durch die Beobachtung der Sterne ermittelten Zeit abwichen. Nachdem alle möglichen Fehlerquellen ausgeschlossen waren, blieb nur eine Erklärung übrig: Erstmals war experimentell nachgewiesen worden, dass sich die Erde im Laufe eines Jahres nicht absolut gleichmäßig dreht, sondern jahreszeitlichen Schwankungen ausgesetzt ist. Darüber hinaus gab es in unregelmäßigen Abständen unvorhersehbare Abweichungen von der erwarteten Zeit, die wohl von anderen Einflüssen auf die Erddrehung herrührten.

Diese Entdeckung der unregelmäßigen Erddrehung durch eine Uhr revolutionierte die Zeitmessung, die sich seit jeher an der scheinbaren Bewegung der Gestirne am Himmel orientiert hatte. Es zeichnete sich ab, dass die Physiker mit elektronischen Uhren die Zeit genauer bestimmen als

Bild 6:
Quarzuhr von Adolf Scheibe und Udo Adelsberger, Physikalisch-Technische Reichsanstalt, Berlin, um 1933. Mit einer dieser Uhren wurde die Unregelmäßigkeit der Erddrehung nachgewiesen. [PTB, Braunschweig]



die Astronomen. Dennoch sollte es noch mehrere Jahrzehnte dauern, bis sich ein Konsens herausbildete, die Zeit nicht mehr von den Sternen abzulesen, sondern in einem von Menschen geschaffenen Apparat aufzubewahren. Aber erst die Atomuhren besiegelten den Wandel von der Zeit der Observatorien zu derjenigen der physikalischen Labore.

Das Sonnensystem als Zeitmaßstab

Spätestens um 1940 war den meisten Astronomen klar: Die Erddrehung war kein absolut zuverlässiger Maßstab zur Bestimmung der Zeit. Deshalb nahm man einen Vorschlag auf, den der Astronom André Danjon aus Straßburg erstmals 1927 geäußert hatte. Die Zeit könnte man zuverlässiger als bisher bestimmen, indem man nicht den Umlauf der Planeten um sich selbst, sondern um die Sonne beobachtet [12]. Auf den Punkt gebracht bedeutete Danjons Idee, dass bei der Zeitmessung anstelle des Tages das Jahr treten sollte.

Dieser Vorschlag hatte einiges für sich: Denn schon seit Jahrhunderten gehörte es zum Allgemeinwissen, dass sich die Planeten in regelmäßigen Bahnen um die Sonne bewegen. Verbreitet wurde dieses Wissen um das von Isaac Newton formulierte Gravitationsgesetz seit dem 18. Jahrhundert unter anderem durch Planetarien. Sie veranschaulichten mit Zahnradgetrieben, wie die unsichtbare Kraft der Gravitation den Weg der Planeten um die Sonne beeinflusst.

Die Astronomen konnten folglich die Position der Sonne und der Planeten bestimmen, wenn sie die Zeit und das Datum wussten. Diese Tabellenwerke für die Positionen der Sonne und der Planeten werden traditionell als „Ephemeriden“ bezeichnet.

Die Ephemeriden konnte man aber auch zur Zeitmessung verwenden. Dabei diente die Sonne als Referenzpunkt. 1948 veröffentlichte Gerald Clemence vom *US Naval Observatory* in Washington D.C. einen detaillierten Vorschlag. Er griff dabei auf Berechnungen der Position der Sonne zurück, die seit etwa 1900 nach einer Formel des Astronomen Simon Newcomb bestimmt worden waren.

1950 stellte Clemence seine Idee vor der *International Astronomical Union* in Paris vor, deren Direktor Danjon damals war. 1952 akzeptierte die Generalversammlung der IAU den Vorschlag von Clemence. Die neue Form der Zeitmessung sollte Ephemeriden-Zeit heißen. Die Sekunde beruhte nun auf der Länge des Jahres 1900, also der Zeitspanne, die die Erde für ihren Umlauf um die Sonne bezogen auf den Fixsternhimmel braucht. Dieses bereits damals recht weit in der Vergangenheit liegende Jahr wurde gewählt, weil die Jahreslänge eben doch nicht gleich lang ist – aber die Veränderungen waren klein und recht konstant, so dass man sie berechnen konnte.

Die Generalversammlung wich in einem entscheidenden Punkt von Clemence ab: Grundlage des Jahres sollte nicht der vollständige Abschluss der Drehung der Erde um die Sonne sein, der als siderisches Jahr bezeichnet wird. Denn durch das Trudeln der Erde um die eigene Achse ist das siderische Jahr um etwas mehr als 20 Minuten länger als das tropische Jahr, das man stattdessen als Bezugsgröße setzte. Als tropisches Jahr bezeichnet man die Spanne zwischen einer Tag- und Nachtgleiche im Frühling zur nächsten. Der Grund für die Wahl des tropischen Jahrs: Hätte man das siderische gewählt, so hätte sich der Termin des Frühjahrsbeginns allmählich verschoben – immerhin etwa alle 70 Jahre um einen Tag.

1956 empfahl die internationale Vereinigung der Astronomen: Als Basiseinheit der Sekunde gelte der 31.556.925,9747te Teil des tropischen Jahres am 0. Januar 1900 [= 31. Dezember 1899] um 12 Uhr Ephemeriden-Zeit. Diese neue Definition der Sekunde wurde 1960 auch von der *Conférence Générale des Poids et Mesures*, dem obersten Gremium für Maße und Gewichte, übernommen und in das neue System von Basiseinheiten integriert.

Als die Ephemeriden-Sekunde 1960 zum internationalen Standard wurde, lief die erste Caesium-Atomuhr von Louis Essen bereits fünf Jahre. Das Ende der von Astronomen anhand der Erddrehung ermittelten Zeit war absehbar. Denn schon unmittelbar nach der Inbetriebsetzung der ersten Caesium-Atomuhr beim *National Physical Laboratory* 1955 wurde darüber diskutiert, wie die Übergabe der Zeitbestimmung aus den Händen der Astronomen in die der Physiker gestaltet werden könne. Doch zunächst musste erst noch ermittelt werden, wie viele Schwingungen ein Caesium-Atom in einer Sekunde ausführt. Dazu benötigte man noch einmal die Hilfe der Astronomen. Mittels raffinierter Konstruktionen wie der Mondkamera, die William Markowitz am *US Naval Observatory* konstruiert hatte, wurde die Atomzeit mit der Ephemeriden-Zeit harmonisiert.

Wer meinen sollte, dass damit die Astronomie bei der Zeitbestimmung ausgedient haben sollte, geht gründlich fehl. Denn die Atomzeitskala ist zwar deutlich regelmäßiger als die mittels Erddrehung ermittelte, aber in unserem Alltag spielt der mittlere Sonnentag und damit der Ausgleich zwischen Atom- und Sonnenzeit weiterhin eine nicht zu unterschätzende Rolle.

Dabei muss bedacht werden, dass die Atomsekunde anhand der Ephemeriden-Zeit bestimmt wurde. Die Definition der SI-Sekunde anhand der Schwingungen des Caesiumatoms beruht folglich mittelbar auf Messungen, die das Jahr 1900 betreffen, und auf Beobachtungsdaten, die teils noch ins 18. und 19. Jahrhundert zurückreichen. Doch in den vergangenen 100 Jahren hat sich die

Erddrehung unter anderem durch den Einfluss der Gezeiten verlangsamt. Dadurch ist der mittlere Tag heute um etwa 2 Millisekunden länger als im Jahr 1900. Diese Unterschiede zwischen der Atomzeit und der astronomischen Zeit summieren sich im Mittel auf eine Sekunde etwa jede 500 Tage. Diese Abweichungen werden vom *International Earth Rotation Service* aufmerksam beobachtet. Sobald die Unterschiede zu groß werden, entscheidet diese Behörde, wann eine Schaltsekunde eingefügt wird. Sollte sich die Erdrotation einmal wieder beschleunigen, so ist denkbar, dass auch eine negative Schaltsekunde notwendig werden könnte.

Die Schaltsekunde ist alles andere als unumstritten. Denn jede neue Schaltsekunde gilt als Fehlerquelle für Computersysteme und die zunehmend vernetzte Welt. Deshalb wird darüber diskutiert, ob es nicht sinnvoller wäre, auf die relativ häufige Einfügung von Sekunden zu verzichten und stattdessen in größeren Abständen einen deutlicheren Sprung zu wagen – etwa mit einer ganzen Schaltstunde im Jahr 2600. Als prominentes Beispiel geht das *Global Positioning System* GPS in diese Richtung voraus: Seit seiner Einführung 1980 ignoriert GPS jede neu hinzu gekommene Schaltsekunde. Doch zunächst wurde die Entscheidung über die Zukunft der Schaltsekunde vertagt. Die *ITU World Radiocommunications Conference* hat 2015 beschlossen, erst einmal an der bewährten Praxis festzuhalten [13]. Wie es in Zukunft mit der (Schalt-)Sekunde weitergeht, berät die Konferenz wieder in 2023.

Literatur

- [1] Otto Neugebauer, „*A history of ancient mathematical astronomy*“, 3 Bände, Berlin, Heidelberg, New York 1975, hier Bd. 1, S. 367
- [2] Wie das folgende; Gerhard Dohrn-van Rossum, „*Die Geschichte der Stunde. Uhren und moderne Zeitordnungen*“, München 1995, S. 46
- [3] Gerhard Dohrn-van Rossum, „*Die Geschichte der Stunde. Uhren und moderne Zeitordnungen*“, München 1995, S. 260f
- [4] Klaus Maurice, „*Die deutsche Räderuhr*“, 2 Bände, München 1976, hier Bd. 1, S. 146, Anm. 87
- [5] Wie das folgende; Klaus Maurice, „*Jost Bürgi oder über die Innovation*“, in; Klaus Maurice, Otto Mayr (Hrsg.), „*Die Welt als Uhr. Deutsche Uhren und Automaten 1550–1650*“, München/Berlin 1980, S. 90–104
- [6] Klaus Maurice, „*Jost Bürgi oder über die Innovation*“, a. a. O., S. 90–104, hier S. 93
- [7] Wie das folgende; Silvio A. Bedini, „*Die mechanische Uhr und die Wissenschaftliche Revolution*“, in; Klaus Maurice, Otto Mayr (Hrsg.), „*Deutsche Uhren und Automaten 1550–1650*“, München/Berlin 1980, S. 21–29
- [8] Zur erst allmählichen Einführung der mittleren Zeit an öffentlichen Uhren vgl.: Jakob Messerli, „*Gleichmäßig, pünktlich, schnell. Zeiteinteilung Zeitgebrauch in der Schweiz im 19. Jahrhundert*“, Zürich 1995, S. 57–68
- [9] Catherine Cardinal, „*La révolution dans la mesure du temps. Calendrier républicain heure décimale 1793–1805*“, La Chaux-de-Fonds 1989
- [10] Zur Geschichte der Shortt-Uhren und ihrer Rolle bei der Entdeckung der unregelmäßigen Erdrotation vgl.; Robert A. Miles, „*Synchronome. Masters of Electrical Timekeeping*“, 2011
- [11] Horst Hassler, „*Adolph Scheibe und Udo Adelsberger. Physiker und Uhrenbauer*“, in; Johannes Graf (Hrsg.), „*Die Quarzrevolution. 75 Jahre Quarzuhr in Deutschland*“, Furtwangen 2008, S. 24–39
- [12] Zur Geschichte der Ephemeriden-Sekunde vgl.; Tony Jones, „*Splitting the second. The story of atomic time*“, Bristol / Philadelphia 2001, S. 17–24
- [13] *Coordinated Universal Time (UTC) to retain “leap second”*, Pressemitteilung der ITU vom 19. November 2015 (https://www.itu.int/net/pressoffice/press_releases/2015/53.aspx; abgerufen am 16.06.2017)

Atomare Definition der Zeiteinheit 1967–2017

Andreas Bauch*

1. Einleitung

«Vous avez déjà dû donner au Monde une nouvelle définition du mètre, basée sur la longueur d'onde d'une radiation lumineuse. Voici que vous allez, au cours de cette session, rechercher dans la profondeur de l'atome l'étalon de définition du temps, celui de la seconde. Vous nous offrez un beau sujet de méditation; la mesure de la course des étoiles dans un cosmos infiniment grand, à l'aide de la vibration d'un atome infiniment petit !»

„Sie haben bereits der Welt eine neue Definition des Meters zur Verfügung gestellt, die auf der Wellenlänge einer Lichtstrahlung beruht. So werden Sie im Verlauf dieser Sitzung in der Tiefe des Atoms das Normal für die Definition der Zeit – die Sekunde – erforschen. Sie bieten uns ein schönes Thema zum Reflektieren: die Messung der Sternbahnen in einem unermesslich weiten Kosmos mithilfe der Schwingung eines unendlich kleinen Atoms!“

Mit diesen Worten eröffnete der französische Außenminister Maurice Couve de Murville die 13. Sitzung der Generalkonferenz für Maß und Gewicht (CGPM, von *Conférence générale de poids et mesures*) im Oktober 1967. In der Tat erfolgte der Beschluss, die Sekunde neu zu definieren, und zwar wie folgt (siehe am Ende von Anhang 1):

«La seconde est la durée de 9 192 631 770 périodes de la radiation correspondant à la transition entre les deux niveaux hyperfins de l'état fondamental de l'atome de césium 133.»

„Die Sekunde ist das 9 192 631 770-fache der Periodendauer der dem Übergang zwischen den beiden Hyperfeinstruktur-niveaus des Grundzustandes von Atomen des Nuklids ^{133}Cs entsprechenden Strahlung“.

In diesem Heft der PTB-Mitteilungen soll in verschiedenen Beiträgen aufgezeigt werden, wie es zu diesem Beschluss, zu dem Zahlenwert und zu

dem Wortlaut kam. Offensichtlich musste zuerst „die Atomuhr“ erfunden werden. Einige Grundlagen und Voraussetzungen werden im nächsten Kapitel behandelt, und der Leser wird bemerken, dass Physiker aus England und USA diese Entwicklung wesentlich befördert haben. Dies behandeln die zwei Beiträge von Mitarbeitern aus den damals aktiven Metrologieinstituten, von Peter Whibberley vom *National Physical Laboratory* (NPL), Teddington, UK, und von Michael Lombardi vom *National Institute of Standards and Technology* (NIST), Gaithersburg bei Washington DC, und Boulder, Colorado, US, damals noch *National Bureau of Standards* (NBS) genannt. Der Leser mag in diesen beiden Artikeln gewisse Überschneidungen – auch mit Inhalten des vorliegenden Artikels – finden. Dies erleichtert aber das Verständnis jedes der eigenständigen Beiträge und wurde bewusst in Kauf genommen. Als Autor im Jahr 2017 benutzt man dankbar die zahlreichen Publikationen, die in den letzten sechzig Jahren entstanden und die viel detaillierter die Entwicklung nachzeichnen, als es hier möglich ist [1–8].

Die Entscheidung über die Neu-Definition einer Einheit ist wissenschaftlich wie wissenschaftspolitisch ein bedeutender Akt. Die Akteure werden in Kapitel 4 vorgestellt und, soweit zugänglich und aus Platzgründen möglich, die Entscheidungsfindung von den Anfängen bis in das Jahr 1967 nachgezeichnet. Es kommen auch die Alternativen zur Caesiumdefinition zu ihrem Recht, und Beiträge aus anderen Ländern werden vorgestellt. Bis zu diesem Zeitpunkt war die Rolle der PTB eher bescheiden und wird in Kapitel 5 skizziert. Soweit aus den Dokumenten zu erschließen ist, trug ein Beitrag aus der PTB zumindest zur Klärung bei, dass mit der neuen Sekundendefinition die Einheit der *Eigenzeit* (*proper time unit*) definiert werden sollte und daher eine explizite Erwähnung der Relativitätstheorie nicht notwendig ist bzw. sogar irreführend wäre. Der Beitrag von Claus Kiefer in diesem Heft wird dabei helfen, diese Thematik zu verstehen. Seit dem Jahr 1967 wurde die Unsicherheit, mit der

* Dr. Andreas Bauch, Fachbereich „Zeit und Frequenz“, E-Mail: andreas.bauch@ptb.de

Caesium-Atomuhren die SI-Sekunde realisieren, um nahezu eine Größenordnung pro Dekade reduziert. In Kapitel 6 skizziere ich knapp die stattgefundene Entwicklung. Die Atomuhrenentwicklung und -verbreitung hat zahlreiche Konsequenzen für technische und wissenschaftliche Anwendungen. Der Beitrag von Gerhard Beutler soll hier stellvertretend über die Nutzung der Atomuhr im Bereich der Geodäsie berichten. Weil es immer gut ist zu wissen, woher man kommt, habe ich Johannes Graf gebeten, den historischen Kontext noch weiter zurückblickend zu beschreiben. Weil es weiterhin keinen Stillstand in der Uhrenentwicklung gibt, sondern im Gegenteil in kurzer zeitlicher Folge immer neue faszinierende Ergebnisse publiziert werden, wird im Beitrag von Ekkehard Peik die Zukunft der Zeiteinheit beleuchtet.

Weiterführendes Material ist in zwei Anhängen zusammengefasst, um den Lesefluss im vorliegenden Artikel nicht zu unterbrechen. Im ersten sind einige Zitate aus Protokollen (*procès verbaux*, *rappports*) der Sitzungen der diversen Komitees in der Originalsprache enthalten.

Im zweiten wird eine Kurzbeschreibung der Funktion der Caesium-Atomuhr gegeben.

2. Die Entwicklung der ersten Atomuhren

Die Wechselwirkung zwischen dem positiv geladenen Kern und der Hülle aus negativ geladenen Elektronen der Atome führt zur Ausbildung von stabilen Konfigurationen der Elektronenhülle eines Atoms. Die Quantenmechanik sagt u. a. voraus, wie die Bindungsenergie einer solchen Konfiguration, auch Eigenzustand des Atoms genannt, von Fundamentalkonstanten, wie dem Elektron-zu-Proton-Massenverhältnis, der Elementarladung, der Lichtgeschwindigkeit und der sogenannten Feinstrukturkonstanten α abhängt. Im Kontext dieses Aufsatzes erscheint die Aussage gerechtfertigt, dass die Bezeichnung „Konstanten“ korrekt ist, der Energieabstand zwischen Eigenzuständen daher ebenfalls konstant ist. Schon 1870 hatte der englische Physiker James Clerk Maxwell vorgeschlagen, grundsätzlich derartige Naturkonstanten für die Festlegung der physikalischen Einheiten zu verwenden und sich nicht auf von der Erde gelieferte Maße, wie die Tageslänge für die Sekunde oder den Erdumfang für das Meter, zu beziehen.

Ein Übergang zwischen zwei atomaren Eigenzuständen mit einer Energiedifferenz ΔE ist mit der Absorption oder Emission elektromagnetischer Strahlung der Frequenz $f = \Delta E/h$ verknüpft. Die Frequenz f bzw. die Periodendauer $1/f$ einer solchen Strahlung ist nach dem Voranstehenden prinzipiell konstant, anders als die Periode der Erdrotation, erst recht aber als die Schwingungsdauer eines Pendels, sei es noch so aufwendig

konstruiert und aufgehängt. Ein bestimmtes Vielfaches von $1/f$ als neue Zeiteinheit zu definieren entsprach der Idee Maxwells und wurde erstmals 1940 von dem amerikanischen Physiker Isidor Isaac Rabi konkret vorgeschlagen: Seine Methode der *molecular beam magnetic resonance* könne man nicht nur zur Untersuchung atomarer Eigenschaften verwenden, schrieb er, sondern, quasi in Umkehrung, die Übergangsfrequenz zwischen zwei ausgewählten atomaren Zuständen sei als Referenz für ein Frequenznormal nutzbar. Er identifizierte die Hyperfeinstrukturzustände im Atom ^{133}Cs als hierfür besonders geeignet [9]. Mit dem Nobelpreis 1944 ausgezeichnet, schaffte es sein Vorschlag als „*radio frequencies in hearts of atoms would be used in most accurate of time-pieces*“ an prominenter Stelle in die Ausgabe der *New York Times* vom 21. Januar 1945.

Rabis Vorschlag hatte einerseits eine Vorschichte und andererseits dauerte es weitere 10 Jahre, ehe die erste Caesium-Atomuhr „tickte“, hierzu noch einige Bemerkungen [8]. Für die Entstehung der Caesium-Atomuhr unerlässlich waren drei Errungenschaften: die Erzeugung von Atomstrahlen im Vakuum, das Verständnis der Richtungsquantisierung, d. h. der Ausrichtung magnetischer Momente von Atomen im Raum und ihrer Manipulationsmöglichkeiten, und die Entwicklung elektronischer Komponenten zur Erzeugung der notwendigen Anregungsstrahlung und deren Kontrolle. Die ersten beiden Errungenschaften sind untrennbar verbunden mit dem Namen Otto Stern, Professor an den Universitäten von Frankfurt und Hamburg (ab 1923) [10]. In seinen ersten Experimenten konnte er die Geschwindigkeitsverteilung von Atomen in einem Atomstrahl ermitteln. Das zusammen mit Walter Gerlach durchgeführte „Stern-Gerlach-Experiment“ begegnet praktisch jedem Physiker in der „Einführung in die Atomphysik“ – in Vorlesung oder Lehrbuch. Es erbrachte den Nachweis, dass das magnetische Moment eines Atoms in einem äußeren Magnetfeld nicht beliebige Orientierungen, sondern diskrete Werte einnimmt. Die Kraft auf ein Atom mit magnetischem Moment μ in einem inhomogenen Magnetfeld nimmt dann ebenfalls diskrete Werte an. Im historischen Experiment wurde die Aufspaltung eines Silber-Atomstrahls nach der Passage eines inhomogenen Magnetfelds beobachtet. Zwei Hamburger Mitarbeiter Sterns, Otto Frisch und Emilio Segrè, blockierten in ihrem Experiment einen der Teilstrahlen hinter dem Magneten „A“ (Polarisator), selektierten den anderen mit einem zweiten Magneten „B“ (Analysator) und induzierten Übergänge zwischen den Zuständen der Richtungsquantisierung in der Zwischenregion „C“ mittels eines statischen Magnetfelds mit schneller Richtungsänderung. Rabi, auch zeitweise Mitarbeiter von Stern und inzwischen an der

Columbia University, bewirkte 1938 im statischen C-Feld durch Radiofrequenz-Einstrahlung bei der Frequenz $f_0 = (E_2 - E_1)/h$ den Übergang zwischen Zuständen der Richtungsquantisierung mit den Energien E_1 und E_2 . Bis heute wird der Begriff „C-Feld“ für den Bereich des schwachen statischen Feldes verwendet, in dem der Hyperfeinstrukturübergang in Caesiumatomen durch die Wechselwirkung mit Mikrowellenstrahlung induziert wird, selbst wenn es in einer Caesium-Fontänenuhr (siehe Kapitel 6) keine Magnete „A“ und „B“ mehr gibt. Die Funktion der Caesium-Atomuhr wird in Anhang 2 beschrieben. Dort wird auch auf eine andere wesentliche Entwicklung aus diesen Jahren, Norman Ramseys Methode der *separated oscillatory fields* [11,12], eingegangen, mit der die Vorteile der Resonanzspektroskopie im Atomstrahl erst voll zur Geltung kommen.

Einen atomaren Übergang anzuregen und nachzuweisen ist sozusagen der erste Schritt. Hierfür ist – für die in der damaligen Zeit diskutierten Atome – durchstimmbare Strahlung im Mikrowellenbereich (Frequenzen 1 GHz bis 30 GHz) notwendig. Diese muss durch Frequenz-Vervielfachung von Radiofrequenz-Signalen im Kilohertzbereich aus einem Quarzoszillator erzeugt werden. Aus der Beobachtung der Wirkung auf die Atome gewinnt man ein Regelsignal zur Steuerung des Quarzoszillators, und durch Frequenzteilung von dessen Signal wiederum Impulse mit der Wiederholrate Eins pro Sekunde. Diese kann man dann zählen und zum Antrieb eines Uhrwerks verwenden. Die Entwicklung der vielfältigen, komplizierten elektronischen Komponenten gelang erst in der Dekade nach dem 2. Weltkrieg, in dessen Verlauf die Radartechnik erhebliche Fortschritte gemacht hatte [4]. Im gleichen Zeitraum wurden auch verschiedene Atome und Moleküle als Kandidaten für eine Atomuhr diskutiert, nämlich Wasserstoff, Rubidium, Caesium, Thallium und daneben das Ammoniakmolekül. Letzteres ist insofern historisch interessant, weil es in der ersten funktionierenden „Atomuhr“ verwendet wurde, wenn auch nicht mit dem gewünschten und erwarteten Gewinn an Genauigkeit, wie Lombardi beschreibt.

Die Entwicklung der ersten Caesium-Atomuhr erfolgte gleichzeitig am amerikanischen NBS (heute NIST) und am britischen NPL. Neben den Beiträgen aus diesen Instituten in diesem Heft sind die persönlichen Erinnerungen von Louis Essen [6] lesenswert. Das NPL hatte die Nase vorn, die erste funktionstüchtige Atomuhr wurde 1955 vorgestellt [13]. Die Autoren schätzten ab, dass die realisierte Übergangsfrequenz um weniger als relativ $2 \cdot 10^{-10}$ von dem Wert abweichen sollte, den man mit ideal ungestörten Atomen erreichen könnte. Eine solche „Unsicherheit“ des primären Normals ist die entscheidende Kenngröße im Ranking verschiedener Normale. Die Terminologie

hat sich allerdings erst in den Jahren und Jahrzehnten danach voll entwickelt. Bereits im darauffolgenden Jahr konnte man die ersten Caesium-Atomuhren kaufen: Bis zum Jahr 1960 produzierte die *National Company* ca. 50 Exemplare ihres „Atomicrons“ [4,5]. Bei Lombardi finden wir die Ansicht des kompletten Geräts, in der PTB ist mit Glück ein „Strahlrohr“ vor der Verschrottung gerettet worden und hängt im Gang des „Zeitlabors“, dem Kopfermann-Bau, siehe Bild 1. In den frühen 1960er-Jahren begann die Firma Hewlett-Packard (HP), Caesium-Atomuhren auf den Markt zu bringen. Die ersten Strahlrohre waren noch etwas unhandlich, etwa 60 cm lang, und wurden von einem Unternehmen aus der Firmengruppe *Varian Associates* in Beverly, Massachusetts, produziert. Aber bereits 1964 kam zunächst das Modell 5060 und wenige Jahre später das Modell 5061 auf den Markt. Beide passten in übliche Elektronikschränke und wurden in beträchtlichen Stückzahlen verkauft. Die PTB beschaffte im Jahr 1967 eine Uhr vom Typ 5060, die jedoch nicht erhalten ist. In der ältesten hier erhaltenen Uhr vom Typ 5061A aus dem Jahr 1969 wurden nacheinander bis 1996 drei Caesiumstrahlrohre benutzt. Das zuletzt benutzte wurde danach aufgeschnitten, und die Uhr dient in der PTB als Anschauungsobjekt (Bild 2).



Bild 1:
Das Strahlrohr des Atomichrons der PTB, welches Mitte der sechziger Jahre benutzt wurde, um die Gänge der Quarzuhren in Braunschweig und in Mainflingen zu kontrollieren. Einige Elemente sind farbig markiert, so die magnetische Abschirmung (gelb), der Detektor (blau) und die Ablenkmodule (grün). Die Abschirmung ist aufgeschnitten und so sieht man die beiden Endstücke des Ramsey-Resonators (hellgrün). Der Atomstrahl verläuft teilweise in einem dünnen Metallrohr (rosa).

Bild 2:
Kommerzielle Atom-
uhr vom Typ 5061A
aus dem Jahr 1996.
Oben sieht man das
aufgeschnittene
Strahlrohr mit dem
Caesiumofen links
und dem Detektor
rechts. Die beiden
Ablenkmagnete
sind durch die Hal-
terung verdeckt, der
Ramsey-Resonator
aus Kupfer ist gut
zu erkennen.



Aus den frühen 1960er-Jahren datieren Arbeiten über Atomuhren, die andere Elemente als das Caesium verwendeten. Im *Massachusetts Institute of Technology* (MIT) entstand Ende der 1950er-Jahre der sogenannte Wasserstoffmaser, in dem die Hyperfeinstrukturfrequenz des Wasserstoffatoms bei 1,4 GHz benutzt wird [14]. Bereits im Jahr 1964 wurde ein kommerzieller Wasserstoffmaser, von *Varian Associates* produziert, in Lausanne auf dem *Congrès de Chronometrie* vorgestellt und die Wasserstoff-Hyperfeinstrukturfrequenz mit Bezug auf eine Hewlett-Packard-Caesium-Atomuhr bestimmt [5]. Vom Wert der Caesium-Hyperfeinstrukturfrequenz wird gleich noch zu reden sein.

Der früheste Beleg der Nutzung der Hyperfeinstrukturfrequenz des Thallium bei 21,3 GHz datiert aus dem Jahr 1962 [15]. In diesen Jahren war das *Laboratoire Suisse de Recherches Horlogères* (LSRH) sehr aktiv und entwickelte parallel dazu auch eine Caesium-Atomuhr à la NPL und NBS. Die höhere Frequenz des Thalliumübergangs sollte à priori ein Vorteil bezüglich erreichbarer Stabilität und Genauigkeit der Uhr sein, erwies sich aber anfangs als große technische Herausforderung. Der langfristig wohl schwerwiegendere Nachteil des Atoms Thallium sind die geringen Kräfte, die mit inhomogenen Magnetfeldern aufgebracht werden können, sodass die erreichbaren Ablenkwinkel klein sind und zur Trennung der beiden Atomzustände eine lange

Apparatur notwendig ist (siehe Anhang A2). Zur Erzeugung eines intensiven Atomstrahls war außerdem eine Temperatur von ca. 700 °C notwendig. An eine Miniaturisierung war nicht zu denken, und nach 1967 findet man praktisch keine Erwähnung dieses Atomuhr-Typs mehr.

In den Diskussionen zur Neudefinition wurde das Für und Wider dieser drei möglichen Kandidaten abgewogen, das Rubidium war schon 1965 „aus dem Rennen“. Das Prinzip der Realisierung der Rubidium-Atomuhr, in [16, 17] beschrieben, führt zu großen systematischen Verschiebungen der Resonanzfrequenz, die schwer zu modellieren sind. Schon früh wurde allerdings die Möglichkeit der Miniaturisierung erkannt und umgesetzt, und heute basiert die mit Abstand größte Zahl produzierter Atomuhren auf Rubidium. Die Signale der Satelliten des *US Global Positioning System* GPS werden heute überwiegend von Rubidium-Atomuhren an Bord abgeleitet, siehe hierzu den Artikel von Gerhard Beutler in diesem Band.

3. Die Bestimmung des Zahlenwerts 9 192 631 770

Die Sekunde als Einheit der Zeit war ursprünglich de facto definiert als Bruchteil des mittleren Sonnentages. Die Internationale Astronomische Union empfahl 1955 stattdessen die Sekunde der Epheme-

riden-Zeit als Zeiteinheit zu verwenden. Denn der jährliche Umlauf der Erde um die Sonne ist in der Tat gleichförmiger und besser vorhersagbar als die Drehung der Erde um ihre Achse. Ausgangspunkt ist das sogenannte tropische Jahr, die Zeitdauer zwischen zwei aufeinanderfolgenden Frühlings-Tag-und-Nacht-Gleichen. Da diese Zeitspanne wegen Präzession und Nutation der Erdachse in bekannter, gesetzmäßiger Weise veränderlich ist, wurde die Sekunde als Bruchteil $1/31\,556\,925,9747$ des differentiellen tropischen Jahrs mit dem 31. Dezember 1899, zwölf Uhr Weltzeit, als Mittelpunkt vorgeschlagen [1]. Die festgelegte Zahl stützt sich auf eine Berechnung von Simon Newcomb aus dem Jahr 1895, der astronomische Beobachtungen aus den Jahren 1750–1892 ausgewertet und das mittlere Verhältnis der Perioden von Erdrotation und Erdumlauf ermittelt hatte. Im Prinzip erfordert die Angabe eines Zeitpunkts in Ephemeriden-Zeit die Bestimmung der Position der Sonne vor dem Fixsternhimmel. Das geht in vernünftiger Zeit nur sehr ungenau. Stattdessen behilft man sich mit der Bestimmung der Mondposition vor den Fixsternen, und hierfür entwickelte William Markowitz vom *United States Naval Observatory* (USNO) die sogenannte *dual-rate moon camera*. Wie genau man damit letztlich die Dauer der Ephemeriden-Zeit bestimmen konnte ist wohl auch unter Astronomen umstritten [1, 18].

Von 1955 bis 1958 wurde in Zusammenarbeit zwischen NPL und USNO die Dauer der Ephemeriden-Sekunde zu $9\,192\,631\,770$ Perioden der Cs-Übergangsfrequenz bestimmt [19]. Als Unsicherheit dieses Zahlenwerts wurde 20 angegeben, eine vermutlich sehr optimistisch abgeschätzte Zahl, die praktisch vollständig durch die astronomische Bestimmung der Ephemeriden-Sekunde bestimmt war. Dessen ungeachtet wurde dieser Zahlenwert von der CGPM in 1964 „for temporary use“ empfohlen* und bildete die Grundlage der 1967 von der 13. CGPM beschlossenen und bis heute gültigen Definition der Zeiteinheit im Internationalen Einheitensystem (SI). Mit den beteiligten Gremien werden wir uns im folgenden Kapitel beschäftigen.

4. Das Faktische bestimmt das Geschehen

Die Sekunde ist die Einheit der Zeit im Internationalen Einheitensystem SI, welches in der Zeit von 1948 bis 1960 von den Organen der Meterkonvention entwickelt und eingerichtet wurde. Letztere wurde 1875 auf Grundlage einer internationalen Vereinbarung von zunächst 17 Staaten geschaffen, mit dem Grundsatz: „*A tous les temps, à tous les peuples*“, also: „Für alle Zeiten, für alle Völker“**. Die Organe sind: Das Internationale Büro für Maß und Gewicht BIPM (von *Bureau international des poids et mesures*), sozusagen das ausführende Organ praktischer Arbeit mit Laboratorien für die wesentlichen

Felder der Metrologie, mit Sitz in Sèvres bei Paris. 18 Delegierte aus verschiedenen Ländern bilden das Internationale Komitee für Maß und Gewicht CIPM (von *Comité international des poids et mesures*), welches jährlich tagt und die Arbeit des BIPM überwacht. Das CIPM macht Empfehlungen an die Generalkonferenz für Maß und Gewicht CGPM (von *Conférence générale des poids et mesures*), bei der sich Delegierte aller Unterzeichnerstaaten der Meterkonvention im Abstand von vier bis sechs Jahren treffen. Das CIPM hat eine Reihe von beratenden Komitees (CCs, von *Comité consultatif*), die sich aus im jeweiligen Feld der Metrologie anerkannten Experten zusammensetzen. Im Jahr 1955 gab es ein solches Komitee im Bereich Zeit und Frequenz noch nicht, und es dauerte noch bis 1988 ehe überhaupt dieses Arbeitsgebiet in das BIPM eingegliedert wurde [7]. Zeit und Frequenz waren für Generationen fest in der Hand der Astronomen, und das Internationale Büro für die Zeit (BIH, von *Bureau international de l'heure*) mit Sitz am *Observatoire de Paris* (OP) kümmerte sich seit 1912 um Erdrotation, Referenzsysteme und in den letzten Jahren bis 1988 auch um Atomzeitskalen [20]. Zwar hatte das CIPM auf seiner Sitzung im Jahr 1956 gerade die aus heutiger Sicht wenig hilfreiche Definition der Ephemeriden-Sekunde als Teil des SI auf den Weg gebracht – die dann erst 1960, also 5 Jahre nach dem Ticken der ersten Atomuhr, von der CGPM sanktioniert wurde. Allerdings wurde auch der Weg für die Schaffung eines „Beratenden Komitees für die Definition der Sekunde“ (CCDS) freigemacht (siehe Anhang 1). Dieses tagte dann erstmals im Juni 1957 am BIPM unter der Leitung seines ersten Präsidenten, des Direktors des OP, A. Danjon***. Es wird naturgemäß über die Uhr des NPL berichtet, über bereits 6 Atomichrons, von denen offenbar eines am NPL und eines in einer französischen Forschungseinrichtung der Telekom betrieben wurde. Auch die anderen Kandidaten für Atomuhren werden vorgestellt. Viel Zeit nehmen zwei mehr grundsätzliche Fragen ein: Worüber reden wir überhaupt? Die Astronomen möchten den Begriff „Uhr“ und auch den Bezug zu Zeit und Zeitskalen einstweilen vermeiden, und man einigt sich auf „atomares Frequenznormal“ (*étalon atomique de fréquence*).

Sollte man den Wert der Übergangsfrequenz des Caesiumübergangs provisorisch festlegen? Die Mehrheit, darunter auch A. Scheibe von der PTB, ist dagegen. „Es gibt so viele Caesiumfrequenzen wie es solcher Art Uhren gibt“ (siehe Anhang 1). Nun, das sollte ja gerade nicht so sein, aber das Atomichron des NPL unterschied sich in der Tat zeitweise von der Uhr NPL-Cs1 um $9 \cdot 10^{-10}$. Eine Empfehlung (*Récommandation*), die an verschiedenen Orten betriebenen Atomuhren über Langwellen- und Längstwellensignale zu vergleichen, wurde verabschiedet [20].

* Siehe hierzu R. Vieweg, 12. CGPM, in Anhang 1

** Das SI-System – und insbesondere seine aktuell diskutierte Weiterentwicklung – waren Gegenstand des Heftes 2/2016 der PTB-Mitteilungen.

*** Alle CCDS-Sitzungsberichte sind zugänglich unter <http://www.bipm.org/en/committees/cc/cctf/publications-cc.html> (Zugriff Juni 2017); auf individuelle Zitate wird daher verzichtet

In den Jahren bis zur zweiten Sitzung des CCDS im April 1961 ist viel geschehen, entsprechend groß ist die Anzahl der Beiträge verschiedener Institute über Fortschritte der Atomuhren-Entwicklung: Caesium, Thallium und Wasserstoff sind Kandidaten, Vergleiche zwischen Atomchrons und Labornormalen haben am NBS und dem NPL stattgefunden, die relativen Abweichungen liegen bei $2 \cdot 10^{-10}$. Zwischen den Normalen NBS-1 und NBS-2 wurde eine Abweichung von weniger als $2 \cdot 10^{-11}$ gefunden, aber es wird festgehalten, dass ein internationaler Vergleich mit dieser Genauigkeit bislang unmöglich ist. Beide Zahlenwerte sind um mindestens eine Größenordnung kleiner als die relative Unsicherheit, mit der man die Dauer der Ephemeriden-Sekunde bei Beobachtungen über wenige Jahre bestimmen kann. Diese Aussagen sind verstreut in den *Procès-Verbaux* der Sitzung, einem Bericht des Schriftführers an das CIPM, aber überwiegend in den Beiträgen der einzelnen Institute. Nur im Bericht des USNO wird die Bestimmung der Caesiumfrequenz zu 9 192 631 770 Hz überhaupt erwähnt, das erscheint etwas verwunderlich.

Andere Diskussionspunkte muten uns heute ebenfalls seltsam an: Ob man sich überhaupt mit Zeitskalen und deren Vergleich befassen sollte, da die CGPM nur ein Mandat für die Befassung mit der fundamentalen Zeiteinheit gegeben habe? Das stelle eine zu enge Sichtweise dar, meinen andere: Jeder, der ein Atomfrequenznormal betreibt, realisiere damit auch eine Zeitskala. Man müsse sich international auf eine einigen und Skalenmaß und -beginn standardisieren. Es dauerte bis 1999, ehe das *Comité consultatif* in «*du temps et des fréquences*» umbenannt wurde – so lange lag der Bezeichnung nach der Fokus auf «*Definition de la seconde*».

Für einen Wortlaut der Neuen Definition gibt es Vorschläge im Annex 4 des Berichts, sie sind im Anhang 1 zu finden.

Das Fazit der Sitzung hatte L. Essen (NPL) gleich vorweggenommen: Atomuhren seien schon heute (1961) besser als astronomische Methoden der Zeitbestimmung – und damit sei es Zeit für eine Neudefinition. Aber welches Atom verwenden?

In der dritten Sitzung des CCDS im Dezember 1963 wurde eine Vielzahl von Argumenten ausgetauscht, ob der rechte Zeitpunkt für eine Neudefinition der Sekunde gekommen sei, falls ja, basierend auf welchem Element und mit welchem Wortlaut. Hätte es zu dem Zeitpunkt schon Einigkeit gegeben, so hätte die im darauffolgenden Jahr stattfindende 12. Sitzung des CGPM vermutlich entsprechend entschieden. Aber zu zahlreich waren noch die Zweifel und die widersprüchlichen Interessen zu bedeutend. In der Tat: Die Definition einer Einheit muss wohlüberlegt sein. Sie spezifiziert einen numerischen Wert, hier die Dauer des fundamentalen Zeitintervalls, und muss als Richtschnur für die praktische Realisierung in einem Normal dienen.

Mit der Realisierung der Einheit in einem Normal kommt man dem idealisierten Konzept nahe, aber wie nahe? Hier kommen die Vor- und Nachteile der verschiedenen Kandidaten ins Spiel. Kennt man die verschiedenen Einflussgrößen auf die mit dem Normal realisierte Zeiteinheit vollständig und kann sie hinreichend genau bestimmen? Die Vertreter des Wasserstoffmasers waren hier sehr optimistisch, aber es waren zu dem Zeitpunkt nur zwei Gruppen tätig, und es gab nur Vergleiche zwischen Masern aus dem gleichen Haus. Dagegen waren die Entwickler von Caesiumuhren anscheinend konservativer, sie konnten aber überwiegend Übereinstimmung verschiedener Normale im Rahmen der Vergleichsmöglichkeiten über den Empfang von Längstwellen- und Langwellensendern berichten. Zusammengefasst gab es die folgenden Argumente:

Änderung sofort:

- Die Caesiumnormale sind eindeutig um zwei Größenordnungen reproduzierbarer und stabiler als mit astronomischen Beobachtungen erreichbar;
- De facto wird die Cs-Sekunde (mit 9 192 631 770 Hz) verwendet, auch wenn das noch nicht offiziell beschlossen war – das geschah erst durch die CGPM 1964;
- Wenn man auf weitere Fortschritte in der Uhrenentwicklung warten will, dann kommt man nie zu einer Entscheidung, denn es wird immer weitere Fortschritte geben.
- Wenn man zu einem späteren Zeitpunkt einen besseren Kandidaten findet, dann kann man seine Frequenz mit Bezug auf die de facto Caesiumfrequenz bestimmen und definieren.

Änderung jetzt noch nicht:

- Man weiß noch zu wenig über Atomuhren;
- Mit der Festlegung auf Caesium zum jetzigen Zeitpunkt demotiviert man Forschung an Thallium und Wasserstoff;
- Solange man noch Zweifel an der besten Eignung des Caesiums hat, sollte man warten.

Zu guter Letzt wird fast einstimmig (bei zwei Gegenstimmen von Astronomen und einer Enthaltung) eine *Récommandation* (S1) verfasst, siehe Annex 1.

Aus den Jahren unmittelbar danach ist rege Kommunikation (per Post) erhalten. Hewlett-Packard bringt seine deutlich verbesserte und praktikablere kommerzielle Uhr auf den Markt,

die dann in so großen Stückzahlen verkauft wird, dass man bald über ein recht gutes statistisches Ensemble verfügt. Uhrentransporte (siehe nächster Abschnitt) belegen die Übereinstimmung der Normale am NBS mit denen in Europa. Neubestimmungen der Caesium-Hyperfeinstrukturfrequenz mit Bezug auf die Ephemeriden-Sekunde ergeben keinen signifikant anderen Wert als den zuvor festgelegten. Trotzdem füllt der Schriftverkehr vor und zum Teil nach der 4. CCDS-Sitzung (12. + 13. Juli 1967) fast einen Leitordner im Zeitlabor der PTB. Es geht nun um den Wortlaut: Sollte man bestimmte Randbedingungen der Realisierung explizit erwähnen (Magnetfeld, Atome in Ruhe, Gravitationsfeld)? Soll die Sekunde in einem bestimmten Gravitationspotential definiert werden? Gerhard Becker aus der PTB hat in seinen schriftlichen Beiträgen vehement für zwei Prinzipien gestritten – und sich durchgesetzt:

Der Definitionstext legt den idealisierten Wert der Caesium-Hyperfeinstrukturfrequenz fest und gilt daher für ungestörte Atome – egal, welche Störungen man kennt bzw. in der Zukunft noch findet.

Die Definition legt die Sekunde der Eigenzeit (*proper time*) fest und gilt daher in jedem Gravitationspotential unabhängig von dessen Wert.

Der Rapport der 4. Sitzung des CCDS aus dem Juli 1967 ist der dünnste aller bis dato vorliegenden. Es war vorab schon alles gesagt. Die verabschiedeten Empfehlungen 1 und 2 wurden dann wörtlich in die Resolution der 13. CGPM im

Jahr 1967 übernommen (siehe Annex 1).

Während der Sitzung wurde Herr Becker gebeten, seine diversen Einlassungen zu der relativistisch korrekten Formulierung noch einmal zusammenzufassen. Er schickte diesen Text dann in französischer Sprache mit Brief vom 21. Juli an Dr. J. Terrien, Direktor des BIPM. Er ist dann als Annexe 4 im Rapport abgedruckt. Darin schlägt er vor, dass für metrologische Anwendungen auf der Erde eine koordinierte Zeitskala, *Temps terrestre* (TT), etabliert werden sollte. Die Internationale Astronomische Union übernahm im Jahr 2000 genau diesen Begriff, und die Internationale Atomzeit TAI ist die Realisierung von TT [7]. Beckers Argumente findet man auch in [21].

5. Die Rolle der PTB

Bis zur Entscheidung zur Neudefinition der Sekunde im Jahr 1967 ist die PTB nicht sonderlich mit eigenen Beiträgen hervorgetreten. Im „Bericht über die Tätigkeiten der PTB im Jahr 1958“ berichten zwei Kollegen, W. Schaffeld und H. Bayer, über den „Aufbau eines Cs-Resonators für genaue Zeitmessung“. Auf der ersten Sitzung des CCDS in 1957 hatte der Vertreter der PTB die Fertigstellung einer Caesium-Atomuhr in der PTB für das Folgejahr angekündigt. Als Eingangsdokument zur 2. Sitzung des CCDS 1961 berichteten Schaffeld und Bayer über Details, ohne dass eindeutig erkennbar wurde, bis zu welchem Stadium die Entwicklung fortgeschritten war. Danach

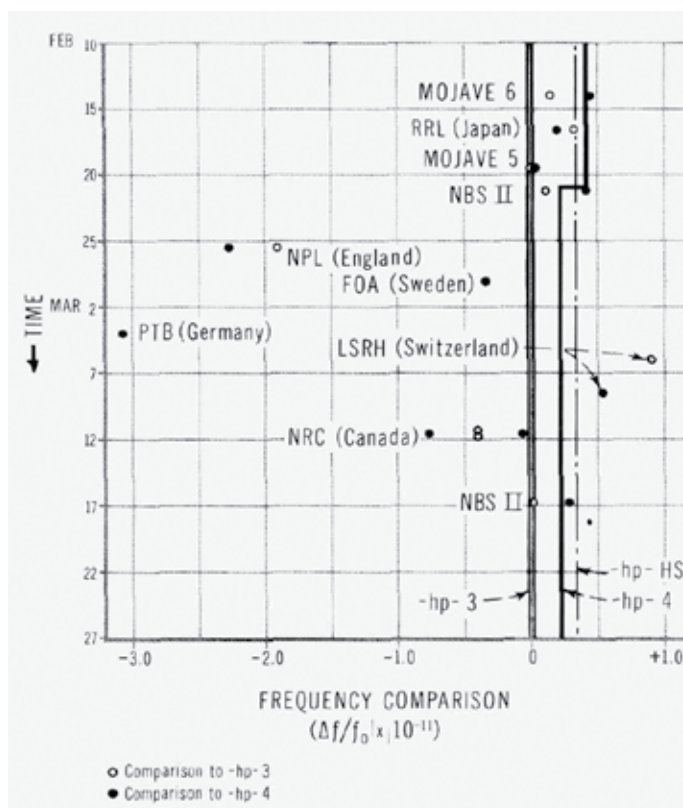


Fig. 9. Frequency comparisons between traveling clocks and other cesium-beam standards are plotted here as frequency differences in parts per 10^{11} . Offset of -150×10^{10} of standards operating on Universal Time has been removed to show only fundamental frequency differences. Results of measurements on all standards shown are well within known accuracies of individual standards.

© Copr. 1949–1998 Hewlett-Packard Co.

Bild 3:
Ergebnis der
Frequenzvergleiche
zwischen zwei rei-
senden Uhren, hp-3
und hp-4 und den
Normalen von NBS,
NPL, LSRH und
dem kanadischen
NRC, das inzwi-
schen auch eine
Eigenentwicklung
vorweisen konnte,
sowie kommer-
ziellen Uhren in
anderen Instituten,
darunter der PTB.
Abbildung aus [22].
Mit Erlaubnis der
Hewlett-Packard
Company.

bestand das Vakuumgefäß, in dem der Atomstrahl fliegen sollte, aus einer 4 Meter langen Glasröhre mit 11 mm Innendurchmesser, die Magnete A und B waren Elektromagnete und der Ramsey-Resonator hatte eine Ausdehnung von 1,1 Meter. Damit erwartete man eine Linienbreite von ca. 120 Hz. Leider verlieren sich damit alle weiteren Spuren dieses Projekts.

Auf der gleichen Sitzung wurde seitens der PTB angekündigt, dass in naher Zukunft die Frequenz der Quarzuhren der PTB mithilfe eines Atomchrons verglichen werden solle. Die ersten Messungen sind für das Jahr 1962 dokumentiert. Wie im vorangehenden Kapitel erwähnt, fanden Vergleiche zwischen den in verschiedenen Einrichtungen betriebenen Atomuhren statt, von denen die ersten von Hewlett-Packard organisiert wurden, nachdem das neue Modell HP-5060 verfügbar war. Zwei Besuche in der PTB sind dokumentiert, für März 1965 [22] und für Juni 1966 [23]. Bild 3 aus [22] zeigt die Ergebnisse von Frequenzvergleichen, darunter die mit der als „90 cm commercial“ bezeichneten Uhr der PTB. Der entsprechende Bericht aus dem Jahr 1966 bezeichnet die Uhr der PTB als „Lab Type“ ohne nähere Angaben, vermutlich handelt es sich hier um einen Irrtum.

Mitte der sechziger Jahre begannen unter Leitung von Gerhard Becker die letztlich erfolgreichen Bemühungen, in der PTB eine eigene

Caesium-Atomuhr zu entwickeln. In der Uhr CS1 (s. Bild 4) setzte man neue Ideen von Holloway und Lacey [24] um. Im Vergleich zu den existierenden Normalen des NPL, des NBS und des LSRH und damit auch zu dem in Anhang 2 gezeigten Prinzip ergab sich eine signifikant andere Konstruktion. Als wesentliche Vorteile wurden erachtet:

- eine reduzierte Frequenzinstabilität durch zweidimensionale Fokussierung der Atome mit magnetischen Linsen (statt der bisher verwendeten Dipolmagnete),
- axialsymmetrische Geometrie des Atomstrahls mit kleiner radialer Ausdehnung,
- Verringerung der Inhomogenität des C-Felds durch Verwendung einer langen Zylinderspule und zylindrischer Abschirmungen – statt eines Magnetfelds quer zur Strahlrichtung.

CS1 wurde aufgebaut und 1969 erstmals benutzt [25] und tickt bis heute. Weltweit sind derzeit nur noch zwei primäre Uhren mit thermischem Atomstrahl in Betrieb, CS1 und CS2 der PTB. CS2 wurde 1985 fertiggestellt, und ihr Konstruktionsprinzip ist dem von CS1 sehr ähnlich. Die



Bild 4:
Die primäre Atomstrahluhr CS1 in der Atomuhrenhalle der PTB (1969).

gegen Ende der 1980er Jahre für CS₂ abgeschätzte Unsicherheit liegt bei $1,2 \cdot 10^{-14}$ [26], also etwa 4 Größenordnungen unter dem Wert der ersten funktionierenden Caesium-Atomuhr. Sie wurde durch Vergleich mit der überlegen genauen Caesium-Fontänenuhr der PTB bestätigt.

6. Die Weiterentwicklung der Caesium-Atomuhr

Jêrome Zacharias vom MIT war der wissenschaftliche „Vater“ des Atomichron, doch er arbeitete bereits in den 1950er-Jahren daneben an einem ambitionierten, doch leider erfolglosen Experiment mit einem aufwärts gerichteten thermischen Atomstrahl. Unter dem Einfluss der Schwerkraft verlangsamten sich in einem solchen Strahl zunächst die Atome, bevor sie nach dem Umkehrpunkt nach unten beschleunigt werden. Zacharias wollte die wenigen besonders langsamen Atome im thermischen Strahl nachweisen, nachdem sie während der Auf- und Abwärtsbewegung mit dem Mikrowellenfeld in Wechselwirkung getreten waren. So hätte sich bei einer mehrere Meter hohen Apparatur eine Linienbreite von unter 1 Hz ergeben, also wesentlich kleiner als die bis dato erreichten ≈ 100 Hz. Die Linienmitte der Resonanzkurve hätte sich so viel leichter auf 0,1 Hz, entsprechend relativ 10^{-11} der Resonanzfrequenz von 9,2 GHz, bestimmen lassen. Allerdings wurde durch Stöße der Atome im Bereich der Ofendüse der winzige, für das Experiment nutzbare, Anteil der extrem langsamen Atome in der thermischen Geschwindigkeitsverteilung noch weiter reduziert. Erhalten geblieben ist der Name „*Fountain*“, den schon Zacharias verwendete.

Ein „*Zacharias fountain*“ konnte erst Realität werden, seit man die Manipulation der Bewegung von Atomen durch Laserstrahlung, speziell die sogenannte Laserkühlung, beherrscht [27]. Durch Laserkühlung lassen sich in einer Zelle mit atomarem Gas Wolken kalter Atome erzeugen, in denen die thermische Bewegung der Atome weitgehend „eingefroren“ ist. Die Temperatur der kalten Atome beträgt dann etwa 1 μ K, und die Verteilung der vorkommenden Relativgeschwindigkeiten ist etwa 10.000-mal schmäler als die thermische Geschwindigkeitsverteilung eines Gases der gleichen Atomsorte bei Raumtemperatur. Es ist ein glücklicher Zufall, dass dies sehr gut mit Caesiumatomen erreicht werden kann. Denn die zur Laserkühlung benötigte schmalbandige und abstimmbare Laserstrahlung bei 852 nm Wellenlänge ist seit gut 20 Jahren mit vergleichsweise einfachen und kompakten Lasern zu erzeugen. Es ist heute Stand der Technik, Fontänenuhren mit vorbestimmter Flughöhe von ca. 1 m der kalten Atome zu realisieren, wobei die räumliche Aufweitung der Wolke mit der Zeit wegen der

schmalen Geschwindigkeitsverteilung gering bleibt. Details findet man in [8, 28]. Die Zeit zwischen den zwei Mikrowellenbestrahlungen beim Durchtritt durch den Mikrowellenresonator im Flug, zunächst aufwärts, dann abwärts, die effektive Wechselwirkungszeit, beträgt in praktisch allen bisher realisierten Fontänenuhren etwa 0,6 s. Die Linienbreite des Uhrenübergangs liegt dann bei 0,8 Hz. Die PTB betreibt derzeit zwei solcher Fontänenuhren, von denen CSF2 im Bild 5 abgebildet ist. Primäre Uhren dieser Art wurden in den Metrologieinstituten von China, Frankreich, Italien, Japan, Russland und den USA entwickelt. Die kleinste abgeschätzte relative Unsicherheit, mit einem solchen Primärnormal die SI-Sekunde zu realisieren, liegt derzeit bei unter $2 \cdot 10^{-16}$.

Bild 5:
Die Caesium-Fontänenuhr der PTB CSF2 (2017)



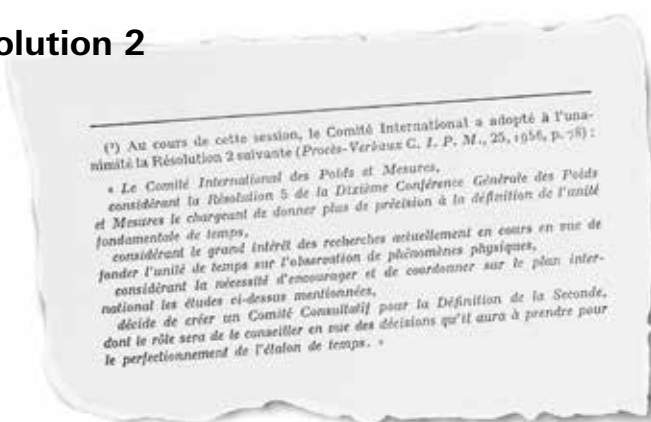
7. Zusammenfassung

Die Caesium-Atomuhr entstand in den 1950er-Jahren und stellte wissenschaftlich wie technisch einen solchen „Quantensprung“ dar, dass Metrologieinstitute, Forschungseinrichtungen, aber auch die Industrie sofort das damit verknüpfte Potenzial erkannten. Zwar gab es andere „Kandidaten“ für eine Neudefinition der Sekunde, aber die im Jahr 1967 getroffene Entscheidung war auch im Rückblick richtig und weitsichtig. Der ursprüngliche Text der Definition war so allgemein formuliert, dass er trotz der Weiterentwicklung der Atom-

uhren über die Jahre von 1955 bis heute nicht in Frage gestellt werden musste. Während dieser Jahre wurde die relative Unsicherheit, mit der die SI-Einheit der Zeit realisiert werden kann, von ca. 10^{-10} bis auf fast 10^{-16} reduziert. Dass es „besser“ geht, wurde gezeigt und ist Gegenstand des Artikels von Ekkehard Peik in diesem Heft. Aber ich sehe derzeit keinen Handlungsbedarf für eine Neudefinition der Sekunde: Für die aktuell gültige gilt weiterhin, dass sie den aktuellen Bedürfnissen entspricht (CGPM 1967, am Ende von Anhang 1).

Anhang 1

CIPM 1956, Résolution 2



„In Anbetracht der Resolution 5, in der die Zehnte Generalkonferenz für Maß und Gewicht (CGPM) das Internationale Komitee für Maß und Gewicht (CIPM) damit beauftragt hat, die Basiseinheit der Zeit mit größerer Genauigkeit zu definieren,

und in Anbetracht der großen Bedeutung der aktuell laufenden Untersuchungen, die das Ziel haben, die Einheit der Zeit auf der Beobachtung physikalischer Phänomene zu begründen,

...

beschließt das Internationale Komitee für Maß und Gewicht (CIPM) die Schaffung eines „Beratenden Komitees für die Definition der Sekunde“. Die Aufgabe dieses Komitees wird es sein, das Internationale Komitee für Maß und Gewicht (CIPM) bei Entscheidungen zur Verbesserung des Zeitnormals zu beraten.“

CCDS_1 (1957)

Sollte man einen Wert für die Caesium-Übergangsfrequenz empfehlen?

„Herr Essen vom NPL bleibt bei seiner Ansicht, dass es eine nutzbringende Aufgabe dieses Komitees wäre, Verwirrungen zu vermeiden, indem ein Zahlenwert festgelegt wird, um die Frequenzaussendungen der verschiedenen Stationen zu koordinieren.

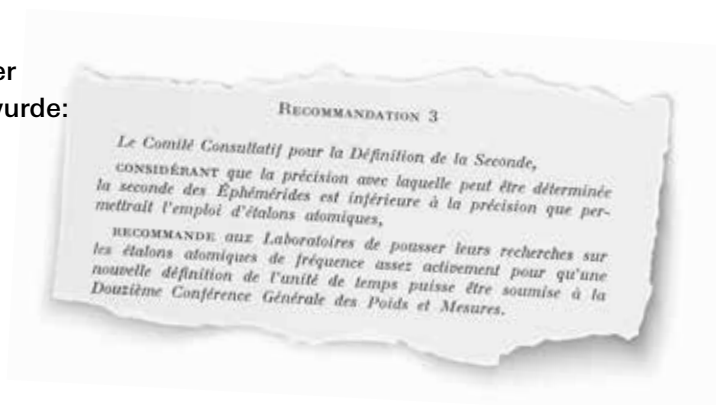
Markowitz [USNO]: Nein, ... es gibt genauso viele Caesiumfrequenzen wie es Caesiumapparaturen gibt.“

« Mr. Essen [NPL] reste d'avis qu'une tâche utile de ce Comité serait d'éviter des confusions, grâce à l'adoption d'une valeur qui servirait à coordonner les émissions de fréquence des diverses stations.

Markowitz [USNO] Non, ... il y a autant de fréquences du césium que d'appareils à césium. »

CCDS-2 (1961)

Eine Empfehlung, die sinngemäß seither in vielen CC Sitzungen verabschiedet wurde:



„Empfehlung 3:

Unter Berücksichtigung der Tatsache, dass die Genauigkeit, mit der die Ephemeriden-Sekunde bestimmt werden kann, geringer ist als die Genauigkeit, die die Verwendung von atomaren Normalen erlauben würde, empfiehlt das „Beratende Komitee für die Definition der Sekunde“ (CCDS)

den Laboratorien, ihre Forschungen zu den atomaren Frequenznormalen aktiv voranzutreiben, damit der Zwölften Generalkonferenz für Maß und Gewicht (CGPM) eine neue Definition der Einheit der Zeit vorgelegt werden kann.

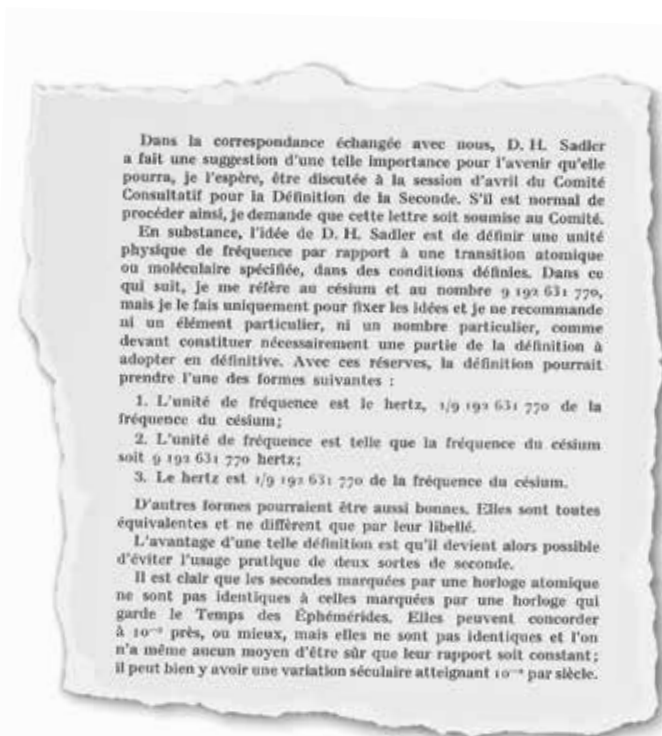
Stellungnahme eines Astronomen vom USNO in dieser Sitzung:

„Im Wesentlichen besteht die Idee darin, eine physikalische Frequenzeinheit mit Bezug auf einen atomaren oder molekularen Übergang unter bestimmten Bedingungen zu definieren. Im Folgenden beziehe ich mich zwar auf Caesium und auf die Zahl 9 192 631 770 – dies allerdings nur, um erst einmal ganz allgemein Ideen festzuhalten. Tatsächlich empfehle ich weder, dass ein bestimmtes Element, noch dass eine bestimmte Zahl in der Definition enthalten sein muss, die letzten Endes festgelegt wird. Berücksichtigt man diese Einschränkungen, so könnte die Definition eine der folgenden Formen annehmen:

1. Die Einheit der Frequenz ist das Hertz, $1/9\,192\,631\,770$ der Caesiumfrequenz.
2. Die Einheit der Frequenz wird so gewählt, dass die Caesiumfrequenz den Wert 9 192 631 770 Hertz hat.
3. Das Hertz beträgt $1/9\,192\,631\,770$ der Caesiumfrequenz.

Andere Formen könnten ebenso gut sein. Sie sind alle gleichwertig und unterscheiden sich nur in ihrem Wortlaut.

...



« En substance, l'idée est de définir une unité physique de fréquence par rapport à une transition atomique ou moléculaire spécifiée, dans des conditions définies. Dans ce qui suit, je me réfère au césium et au nombre 9 192 631 770, mais je le fais uniquement pour fixer les idées et je ne recommande ni un élément particulier, ni un nombre particulier,

...

...

Durch eine solche Definition würde verhindert, dass in der Praxis zwei Arten von Sekunden verwendet werden.

Es ist klar, dass die Sekunden einer Atomuhr nicht identisch mit den Ephemeriden-Sekunden sind. Sie können zwar auf fast 10^{-9} – oder besser – übereinstimmen, aber sie sind nicht identisch und man kann nicht einmal sicher sein, dass ihr Verhältnis konstant ist; ein Unterschied von bis zu 10^{-8} im Verlauf von hundert Jahren ist durchaus möglich.

Ebenso klar ist, dass die beiden Uhrentypen auch weiterhin verwendet werden. Der Astronom benötigt die Ephemeriden-Sekunde, auf der die dynamische Astronomie beruht. Der Physiker hingegen benötigt die – leicht für das Labor zugängliche – Atomuhr. Die Gesetzgebung kann keine dieser Tatsachen verändern.

Aber Astronomen und Physiker sind keine unterschiedlichen Spezies und leben nicht auf unterschiedlichen Planeten. Die Handlungen der einen beeinflussen die Handlungen der anderen, und manchmal sind Physiker und Astronomen sogar in einer einzigen Person vereint. Doch würde die Verwendung eines einzigen Wortes für zwei grundsätzlich unterschiedliche physikalische Konzepte mit Sicherheit zu Verwirrung und falschen Interpretationen führen.

Für den Astronomen ist das Grundkonzept bei der Bestimmung der Zeit die Epoche (der Zeitpunkt einer Beobachtung) – er erhält die Frequenz durch Ableitung. Für den Physiker ist das Grundkonzept die Frequenz – er erhält die Zeit durch Integration, unter Verwendung der vom Astronomen gelieferten Konstanten (dem Anfangswert der Integration).

Es ist daher völlig logisch und mit den Grundbedürfnissen der beiden Wissenschaften konform, die Einheit der Zeit mithilfe astronomischer Abläufe und die Einheit der Frequenz mithilfe physikalischer Abläufe zu definieren. Dadurch geht nichts verloren, sondern man gewinnt dadurch sogar an Eindeutigkeit sowie an Kürze bei der Beschreibung und an sprachlicher Genauigkeit.“

...

comme devant constituer nécessairement une partie de la définition à adopter en définitive. Avec ces réserves, la définition pourrait prendre l'une des formes suivantes :

L'unité de fréquence est le hertz, $1/9\,192\,631\,770$ de la fréquence du césium ;

L'unité de fréquence est telle que la fréquence du césium soit $9\,192\,631\,770$ hertz ;

Le hertz est $1/9\,192\,631\,770$ de la fréquence du césium.

D'autres formes pourraient être aussi bonnes. Elles sont toutes équivalentes et ne diffèrent que par leur libellé. L'avantage d'une telle définition est qu'il devient alors possible d'éviter l'usage pratique de deux sortes de seconde.

Il est clair que les secondes marquées par une horloge atomique ne sont pas identiques à celles marquées par une horloge qui garde le Temps des Éphémérides. Elles peuvent concorder à 10^{-9} près, ou mieux, mais elles ne sont pas identiques et l'on n'a même aucun moyen d'être sûr que leur rapport soit constant ; il peut bien y avoir une variation séculaire atteignant 10^{-8} par siècle.

Il est clair également que les deux sortes d'horloges continueront d'être utilisées. La seconde de Temps des Éphémérides sur laquelle repose l'astronomie dynamique est nécessaire à l'astronome. L'horloge atomique, facilement disponible au laboratoire, est nécessaire au physicien. La législation ne peut modifier aucun de ces faits.

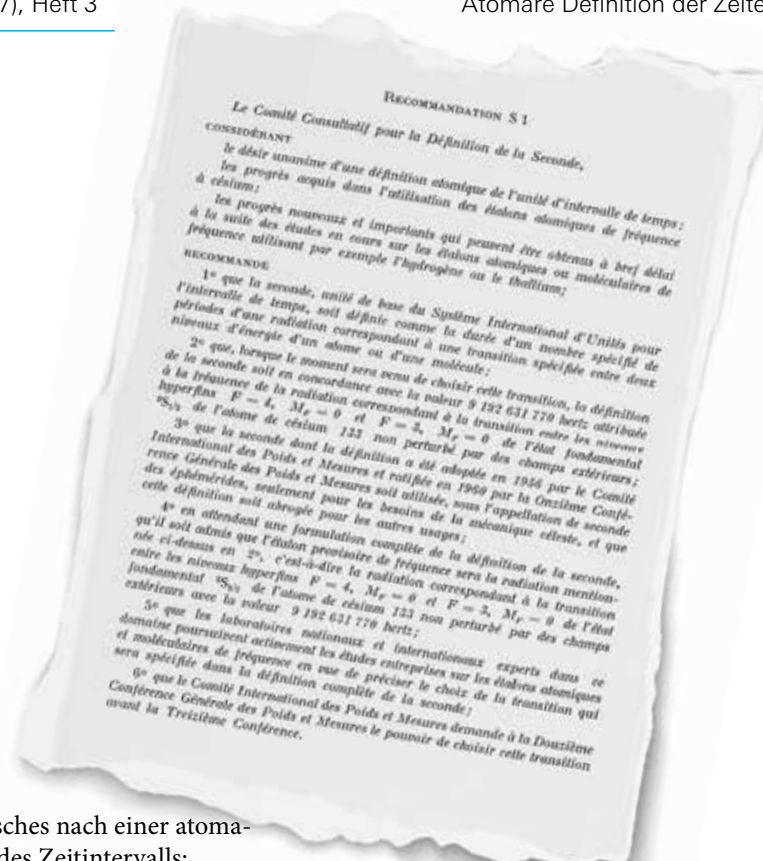
Mais les astronomes et les physiciens ne sont pas deux races distinctes habitant des planètes différentes. Les actions des uns influent sur celles des autres et, parfois, le physicien et l'astronome sont réunis en une seule personne. Il est certain que l'emploi d'un seul mot pour désigner deux concepts physiques fondamentalement différents doit produire des confusions et de fausses interprétations.

Pour l'astronome, le concept fondamental dans la conservation du temps est l'époque ; il obtient la fréquence par dérivation. Pour le physicien, le concept fondamental est la fréquence ; il obtient le temps par intégration, en utilisant des constantes fournies par l'astronome.

Il est donc tout à fait logique, et en accord avec les besoins fondamentaux des deux sciences, de définir l'unité de temps par des opérations propres à l'astronomie et l'unité de fréquence par des opérations propres à la physique. Par ce moyen rien n'est perdu et ce que l'on gagne c'est l'absence de risque de confusion, l'économie dans la description et la précision du langage. »

CCDS-3 (1963)

Recommendation S 1



„Unter Berücksichtigung

- des einhelligen Wunsches nach einer atomaren Definition der Einheit des Zeitintervalls;
- der erzielten Erfolge bei der Verwendung der atomaren Cesiumfrequenznormale,
- der neuen und bedeutenden Fortschritte, die innerhalb kurzer Zeit als Ergebnis laufender Untersuchungen an anderen atomaren Normalen (Wasserstoff, Thallium) gemacht werden können, empfiehlt das „Beratende Komitee für die Definition der Sekunde“ (CCDS),

1. die Sekunde (die Basiseinheit des Internationalen Einheitensystems (SI) für das Zeitintervall) als die Dauer einer bestimmten Periodenanzahl einer Strahlung zu definieren, die einem bestimmten Übergang zwischen zwei Energieniveaus eines Atoms oder Moleküls entspricht,
2. dass – wenn der Zeitpunkt gekommen ist, diesen Übergang zu wählen – die Definition der Sekunde mit dem Wert 9 192 631 770 Hertz übereinstimmen soll, der der Frequenz der Strahlung zugeordnet ist, welche dem Übergang zwischen den Hyperfeinstruktur-niveaus $F=4$, $M_F=0$ und $F=3$, $M_F=0$ des Grundzustands $^2S_{1/2}$ des Atoms 133-Caesium ohne Störung durch äußere Felder entspricht.
3. ...
4. dass die oben genannte Strahlung provisorisch als Frequenznormal zugelassen wird, bis eine endgültige Definition der Sekunde erfolgt ist,
5. ...
6. dass das Internationale Komitee für Maß und Gewicht (CIPM) die Zwölfte Generalkonferenz für Maß und Gewicht (CGPM) um Befugnis bittet, diesen Übergang zu wählen, bevor die Dreizehnte Generalkonferenz für Maß und Gewicht (CGPM) stattfindet.“

« Le CCDS, considérant

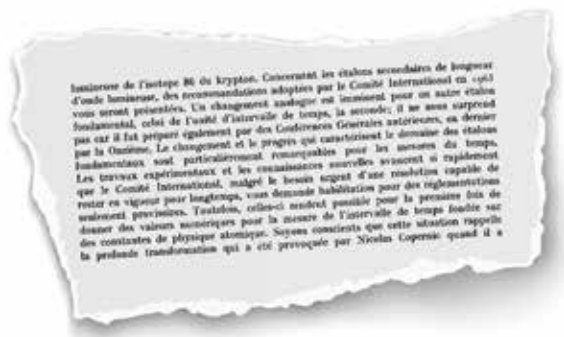
- le désir unanime d'une définition atomique de l'unité d'intervalle de temps ;
- les progrès acquis dans l'utilisation des étalons atomiques de fréquence à césium,
- les progrès nouveaux et importants qui peuvent être obtenus à bref délai à la suite des études en cours sur les étalons atomiques .. (hydrogène, thallium)

Recommande

1. que la seconde, unité de base du Système International d'Unités pour l'intervalle de temps, soit définie comme la durée d'un nombre spécifié de périodes d'une radiation correspondant à une transition spécifiée entre deux niveaux d'énergie d'un atome ou d'une molécule,
2. - que, lorsque le moment sera venu de choisir cette transition, la définition de la seconde soit en concordance avec la valeur 9 192 631 770 hertz attribuée à la fréquence de la radiation correspondant à la transition entre les niveaux hyperfins $F=4$, $M_F=0$ et $F=3$, $M_F=0$ de l'état fondamental $^2S_{1/2}$ de l'atome de césium 133 non perturbé par des champs extérieurs;
3. ...
4. en attendant une formulation complète de la définition de la seconde, qu'il soit admis que l'étalon provisoire de fréquence sera la radiation mentionnée ci-dessus...
5. ...
6. que le CIPM demande à la 12. CGPM le pouvoir de choisir cette transition avant la 13. CGPM. »

12. CGPM Sitzung : 6.–13. Oktober 1964

aus der Eröffnungsrede von Richard Vieweg, Präsident der PTB und des CIPM :



„Die im Bereich der Basisnormale erzielten Änderungen und Fortschritte sind insbesondere bei der Zeitmessung bemerkenswert. Die experimentellen Arbeiten und die neuen Erkenntnisse schreiten so schnell voran, dass das Internationale Komitee Sie bittet, trotz des dringenden Bedarfs nach einer Resolution, die von langfristiger Dauer sein wird, vorab schon einmal provisorische Regelungen treffen zu dürfen. Dadurch würde es zum ersten Mal möglich, Zahlenwerte für die Messung der Dauer von Zeitintervallen anzugeben, welche auf den Konstanten der Atomphysik beruhen. Wir sollten uns bewusst sein, dass diese Situation an die tiefgreifende Umwandlung erinnert, die von Nikolaus Kopernikus hervorgerufen wurde, als er das klassische Weltbild mit der Erde als Zentrum durch sein neues, heliozentrisches System ersetzte. Heute sind wir im Begriff, die Zeit nicht nur durch die Bewegung der Planeten im Sonnensystem messen zu können, sondern auch durch die Bewegung der Elektronen in einem atomaren System.“

Le changement et le progrès qui caractérisent le domaine des étalons fondamentaux sont particulièrement remarquables pour les mesures du temps. Les travaux expérimentaux et les connaissances nouvelles avancent si rapidement que le Comité International, malgré le besoin urgent d'une résolution capable de rester en vigueur pour longtemps, vous demande habilitation pour des réglementations seulement provisoires. Toutefois, celles-ci rendent possible pour la première fois de donner des valeurs numériques pour la mesure de l'intervalle de temps fondée sur des constantes de physique atomique. Soyons conscients que cette situation rappelle la profonde transformation qui a été provoquée par Nicolas Copernic quand il remplacé la mécanique céleste classique, avec la Terre comme centre, par son nouveau Système Héliocentrique. Aujourd'hui, nous sommes en train d'adjoindre aux mesures du temps par le mouvement des planètes dans le système solaire, des mesures du temps par le mouvement des électrons dans le système atomique.

13. Sitzung der CGPM, Oktober 1967

RÉSOLUTIONS ADOPTÉES PAR LA 13^e CONFÉRENCE GÉNÉRALE

Système International d'Unités (SI)

Unité de temps (seconde)

RÉSOLUTION

La Treizième Conférence Générale des Poids et Mesures,

CONSIDÉRANT

que la définition de la seconde décidée par le Comité International des Poids et Mesures à sa session de 1956 (Résolution 1) et ratifiée par la Résolution 9 de la Onzième Conférence Générale (1960), puis maintenue par la Résolution 5 de la Douzième Conférence Générale (1964) ne suffit pas aux besoins actuels de la métrologie,

qu'à sa session de 1964 le Comité International des Poids et Mesures, habilité par la Résolution 5 de la Douzième Conférence Générale (1964), a désigné pour répondre à ces besoins un étalon atomique de fréquence à césium à employer temporairement,

que cet étalon de fréquence est maintenant suffisamment éprouvé et suffisamment précis pour servir à une définition de la seconde répondant aux besoins actuels,

que le moment est venu de remplacer la définition actuellement en vigueur de l'unité de temps du Système International d'Unités par une définition atomique fondée sur cet étalon,

DÉCIDE

1^o L'unité de temps du Système International d'Unités est la seconde définie dans les termes suivants :

« La seconde est la durée de 9 192 631 770 périodes de la radiation correspondant à la transition entre les deux niveaux hyperfins de l'état fondamental de l'atome de césium 133 ».

2^o La Résolution 1 adoptée par le Comité International des Poids et Mesures à sa session de 1956 et la Résolution 9 de la Onzième Conférence Générale des Poids et Mesures sont abrogées.

RÉSOLUTION 2

La Treizième Conférence Générale des Poids et Mesures,

CONSIDÉRANT que l'étalon de fréquence à césium est encore perfectible et que des expériences en cours autorisent l'espoir de réaliser d'autres étalons ayant des qualités encore meilleures pour servir à définir la seconde,

INVITE les organisations et laboratoires experts dans le domaine des étalons atomiques de fréquence à poursuivre activement leurs études.

**„Die Dreizehnte Generalkonferenz
für Maß und Gewicht,
erwägt,**

dass die Definition der Sekunde, die vom Internationalen Komitee für Maß und Gewicht bei seiner Tagung von 1956 beschlossen (Resolution 1) und durch die Resolution 9 der Elften Generalkonferenz (1960) ratifiziert, dann durch die Resolution 5 der Zwölften Generalkonferenz (1964) beibehalten worden ist, den derzeitigen Erfordernissen der Metrologie nicht mehr genügt,

dass in seiner Sitzungsperiode von 1964 das Internationale Komitee für Maß und Gewicht, ermächtigt durch die Resolution 5 der Zwölften Generalkonferenz (1964), um diesen Erfordernissen Rechnung zu tragen, ein atomares Caesiumfrequenznormal zur vorläufigen Verwendung empfohlen hat,

dass dieses Frequenznormal jetzt ausreichend erprobt und als ausreichend genau anzusehen ist, um für eine Definition der Sekunde, die den derzeitigen Erfordernissen entspricht, zu dienen,

dass der Augenblick gekommen ist, die zur Zeit gültige Definition der Einheit der Zeit des Internationalen Einheitensystems durch eine auf diesem Normal beruhende atomare Definition zu ersetzen,

entscheidet,

1. Die Einheit der Zeit des Internationalen Einheitensystems ist die mit folgendem Wortlaut definierte Sekunde:
„Die Sekunde ist das 9 192 631 770-fache der Periodendauer der dem Übergang zwischen den beiden Hyperfeinstrukturniveaus des Grundzustandes von Atomen des Nuklids ^{133}Cs entsprechenden Strahlung.“
2. Die vom Internationalen Komitee für Maß und Gewicht bei seiner Tagung von 1956 angenommene Resolution 1 und die Resolution 9 der Elften Generalkonferenz für Maß und Gewicht sind aufgehoben.

Resolution 2

In Anbetracht der Tatsache, dass das Caesiumfrequenznormal noch vervollkommnungsfähig ist und laufende Experimente zur Hoffnung Anlass geben, dass weitere – für die Definition der Sekunde noch besser geeignete – Normale realisiert werden können,

ermuntert die Dreizehnte Generalkonferenz für Maß und Gewicht die Organisationen und Laboratorien, die Expertise auf dem Gebiet der Atomfrequenznormale haben, ihre Studien weiterhin aktiv zu betreiben.

ATEMBERAUBEND.

Ultrapräzise Positioniersysteme
auch für den Einsatz in Vakuum und Tieftemperatur.



PI

MOTION CONTROL
www.pi.ws

Anhang 2

Die Funktion der Atomuhr, gezeigt am Beispiel der Caesium-Atomuhr

Atomare Übergänge, die in den 1950er Jahren für die Nutzung in einer Atomuhr diskutiert wurden, sollten folgende Eigenschaften haben:

- eine Übergangsfrequenz im Mikrowellenbereich, möglichst hoch, aber < 30 GHz
- Unempfindlichkeit gegenüber äußeren Störungen, z. B. durch elektrische und magnetische Felder
- Beobachtbarkeit des Übergangs mit langer Wechselwirkungszeit und einem hohen Signal-zu-Untergrund-Verhältnis.

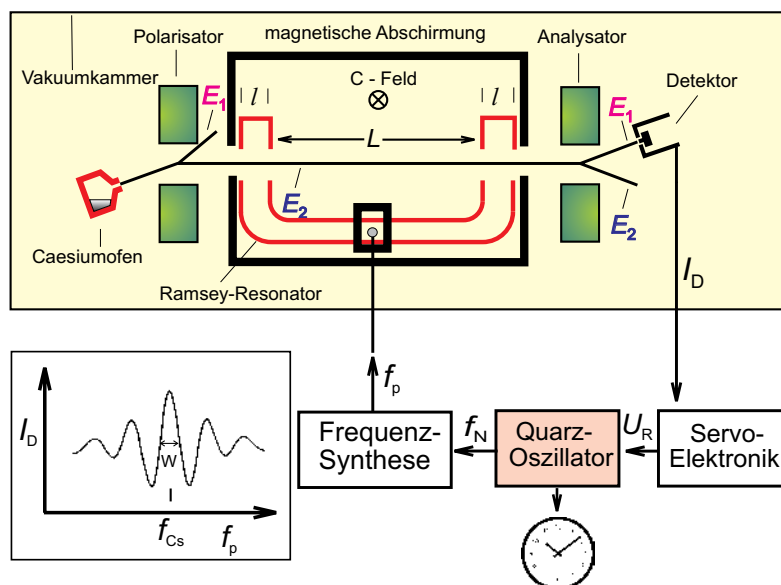
Bild A2:
Cs-Uhr, schematische Darstellung;
 f_N : Normalfrequenz,
 f_p : Frequenz zur Bestrahlung der Atome, I_D : Detektorsignal, U_R : Signal zur Regelung des Quarzoszillators, links unten Detektorsignal I_D als Funktion der Frequenz des Mikrowellenfeldes f_p ; eingezeichnet ist die Linienbreite W .

Den vier Kandidaten, Wasserstoff (H), Rubidium (Rb), Caesium (Cs) und Thallium (Tl) gemeinsam ist die Elektronenkonfiguration: Im Grundzustand kommt es zur Kopplung zwischen dem Spin des einzelnen Leuchtelektrons in der Elektronenhülle des Atoms und dem halbzahligen Kernspin. Diese Kopplung führt zur sog. Hyperfeinstrukturaufspaltung des Grundzustands mit ganzzahligem Spin. Die beiden Untersubzustände mit der magnetischen Quantenzahl $m_F = 0$ definieren den sog. Uhrenübergang, dessen Übergangsfrequenz in erster Ordnung vom statischen Magnetfeld nicht beeinflusst wird. Die Über-

gangsfrequenzen liegen für die vier Elemente bei 1,42 GHz, 6,83 GHz, 9,19 GHz und 21,31 GHz, in der obigen Reihenfolge. Nachfolgend wird die Funktion der Caesium-Atomuhr anhand von Bild A2.1 beschrieben, das im Prinzip auch für die Thallium-Atomuhr gilt. Für Atomuhren mit Wasserstoff und Rubidium sind andere Konstruktionsarten notwendig, die in [2, 14] bzw. [16, 17] beschrieben werden.

Der Uhrenübergang zwischen den in Bild A2 mit E_2 und E_1 bezeichneten Zuständen soll nachgewiesen werden. In der Vakuumkammer einer Atomuhr werden Caesiumatome verdampft, und es wird ein Atomstrahl erzeugt. In diesem kommen zunächst die beiden Zustände gleich häufig vor. Der Magnet (Polarisator) lenkt die Atome so ab, dass nur Atome im Zustand E_2 in den U-förmigen Ramsey-Resonator (s. u.) gelangen. Hier werden die Atome durch Bestrahlung mit einem Mikrowellenfeld der Frequenz f_p in den anderen Zustand E_1 angeregt. Durch den zweiten Magneten (Analysator) werden dann nur die Atome, die eine Zustandsänderung von E_2 nach E_1 erfahren haben, auf den Auffänger gelenkt. Die Anzahl der Atome im Auffänger ist am größten, wenn f_p den für das Caesium-Atom charakteristischen Wert f_{Cs} hat.

Wird die Anregungsfrequenz um die Frequenz des Uhrenübergangs herum variiert, so registriert man eine Resonanzlinie, die in der Abbildung links unten skizziert ist. Ihre spektrale Breite hängt wegen der Abwesenheit von spontanen Übergängen nur von der Dauer der Wechselwirkung mit dem Hochfrequenzfeld ab. Zusätzlich müssen dafür die technischen Voraussetzungen erfüllt sein, dass das statische Feld „C“ über die gesamte Ausdehnung hinreichend homogen und die Phase des anregenden Feldes hinreichend konstant ist. Letzteres ist bei einer Wellenlänge von ca. 3 cm nur über einen sehr kleinen Bereich zu erreichen. Daher fruchtete Norman Ramseys Idee: Das Hochfrequenzfeld wird in einem zweiarmigen Mikrowellenresonator geführt. Die Atome (Geschwindigkeit v) werden dann durch eine Bestrahlung über die Länge von ca. 2 cm im ersten Arm des Resonators in einen kohärenten Superpositionszustand gebracht und nach der freien Driftstrecke der Länge L im zweiten



Resonator-Arm erneut mit dem Hochfrequenzfeld bestrahlt. Die beobachtete Resonanz der Übergangswahrscheinlichkeit hat dann die spektrale Breite $W \approx \nu/(2L)$. Wechselwirkungsstrecken bis zu vier Metern wurden realisiert, in den Uhren der PTB beschränkte man sich aber auf 0,8 m und in den aktuell gefertigten kommerziellen Uhren ist L ca. 20 cm (siehe auch Bild 2 im Text). Bedingt durch die Geschwindigkeit der Atome in einem thermischen Strahl (≈ 200 m/s) ergeben sich damit Linienbreiten zwischen 25 Hz und 500 Hz. Eine elektronische Regelung sorgt dafür, dass der Frequenzgenerator auf der Frequenz f_{Cs} gehalten wird. Diese funktioniert umso besser, je größer der Wert des Quotienten Signalstärke durch Linienbreite ist. Durch Abzählen von 9 192 631 770 Perioden der Mikrowellenstrahlung gewinnt man aus dem Generatorsignal Sekundenpulse, die mit einem Uhrwerk gezählt werden.

Literatur

- [1] G. Becker, „Von der astronomischen zur atomphysikalischen Definition der Sekunde“, PTB-Mitt. **76** (1966), S. 315–323 und S. 415–419
- [2] R. E. Beehler, „A historical review of atomic frequency standards“, NBS Monograph **140** Chapter 4, Boulder (1974)
- [3] R. E. Beehler, R. C. Mockler and J. M. Richardson, „Cesium Beam Atomic Time and Frequency Standards“, Metrologia **1** (1965), S. 114–131 (mit bereits damals 114 Referenzen)
- [4] P. Forman, „Atomichron: the atomic clock from concept to commercial product“, Proc. IEEE **73** (1985), S. 1181–1204
- [5] L. Cutler, „Fifty years of commercial caesium clocks“, Metrologia **42** (2005), S. S90–S99
- [6] National Physical Laboratory, Teddington, UK, and Ray Essen, „The Memoirs of Louis Essen, father of atomic time“ (2015)
- [7] B. Guinot and F. Arias, „Atomic time-keeping from 1955 to the Present“, Metrologia **42** (2005), S. S20–S30
- [8] A. Bauch, E. Peik und S. Weyers, „Wie tickt eine Atomuhr? – Realisierung der Sekunde von 1955 bis heute“, PTB-Mitt. **126** (2016), S. 17–34
- [9] Ref 19 in [4]
- [10] H. Schmidt-Böcking, K. Reich, „Otto Stern, Physiker, Querdenker, Nobelpreisträger“, Biographienreihe der Goethe-Universität Frankfurt, Societäts-Verlag, 2011
- [11] N. F. Ramsey, „A molecular beam resonance method with separated oscillatory fields“, Phys. Rev. **73** (1950), S. 695–698.
- [12] N. F. Ramsey, „Molecular Beams“, London, New York, Oxford University Press (1956)
- [13] L. Essen und J. Parry, „Atomic standard of time and frequency“, Nature **176** (1955), S. 280–284
- [14] H. M. Goldenberg, D. Kleppner, and N. F. Ramsey, „Atomic hydrogen maser“, Phys. Rev. Lett. **5** (1960), S. 361–64
- [15] J. Bonanomi, „A thallium beam frequency standard“, IRE Trans. on Instr. **I–11** (1962), S. 212–215
- [16] C. Audoin and B. Guinot, The Measurement of Time, Cambridge University Press, Cambridge, 2001
- [17] M. E. Packard und B. E. Swartz, „The optically pumped rubidium vapor frequency standard“, IRE Trans. on Instr. **I–11** (1962), S. 215–223
- [18] J. Kovalevski, „Astronomical Time“, Metrologia **1** (1965), S. 169–180
- [19] W. Markowitz, R. G. Hall, L. Essen, J.V.L. Parry, „Frequency of cesium in terms of ephemeris time“, Phys. Rev. Lett. **1** (1958), S. 105–107
- [20] B. Guinot, „History of the Bureau International de l'Heure“, in Polar Motion: Historical and Scientific Problems, ASP Conference Series, Vol. 208, 2000, S. Dick, D. McCarthy and B. Luzum, Hrsg., S. 175–184
- [21] G. Becker, B. Fischer, G. Kramer, und E. K. Müller, „Die Definition der Sekunde und die Allgemeine Relativitätstheorie“, PTB-Mitt. **77** (1967), S. 111–116
- [22] „Correlating time from Europe to Asia with flying clocks“, Hewlett-Packard Journal **16** No. 8, April 1965
- [23] „World-wide time synchronization, 1966“, Hewlett-Packard Journal **17**, No. 12, August 1966
- [24] J. H. Holloway, R. F. Lacey, „Factors which limit the accuracy of cesium atomic beam frequency standards“, Proc. Intern. Conf. Chronometry (CIC 64) (1964), S. 317–331
- [25] G. Becker, B. Fischer, G. Kramer, E. K. Müller, „Neuentwicklung einer Cesiumstrahlapparatur als primäres Zeit- und Frequenznormal an der PTB“, PTB-Mitteilungen **79** (1969), S. 77–80
- [26] A. Bauch, „The PTB primary clocks CS1 and CS2“, Metrologia **42** (2005), S43–S54
- [27] S. Chu, „The manipulation of neutral particles“, Rev. Mod. Phys. **50** (1998), S. 685–706, C. Cohen-Tannoudji, „Manipulating atoms with photons“, ibid. S. 707–720, W. D. Phillips, „Laser cooling and trapping of neutral atoms“, ibid. S. 721–733
- [28] R. Wynands and S. Weyers, „Atomic fountain clocks“, Metrologia **42** (2005), S64–S79

NPL's Contribution to the Introduction of the Caesium Second

Peter Whibberley*

Introduction

For centuries timekeeping was the preserve of astronomers. The fundamental unit of time measurement was the day, sub-divided by clocks for everyday use into hours, minutes and seconds. By the early 20th century, clocks were sufficiently stable to indicate that the length of the mean solar day varied, though their lack of intrinsic accuracy (the rate of a clock depended on its mechanical properties) ensured that the Earth's rotation remained the global reference standard for timekeeping.

This situation changed fundamentally in June 1955, when Louis Essen and Jack Parry brought the first caesium atomic clock into operation at the National Physical Laboratory, in Teddington, UK. The atomic clock not only provided a much more stable timekeeper than the Earth's rotation; its reference was the frequency of an atomic transition - a fundamental constant of nature, determined by the laws of quantum mechanics. As a result all caesium clocks will run at essentially the same rate, limited only by noise processes and their local environment, regardless of time or place.

The development of the atomic clock opened up the prospect of an entirely new way of measuring time, independently of the fluctuating rotation of the Earth. Over the following 12 years, Essen and his colleagues at NPL improved the caesium clock design, investigated methods for performing international time comparisons, and played a leading role in introducing a coordinated global atomic time scale. These were all significant stepping stones towards the redefinition of the SI second in terms of the caesium transition in 1967. This article briefly describes the contributions made by NPL towards the adoption of the atomic second during this dynamic and innovative period.

Early work at NPL on frequency standards

From its foundation in 1900 as the national test, measurement and standardisation laboratory for Britain, the National Physical Laboratory (NPL)

has carried out work in the time and frequency field. Testing of watches and marine chronometers continued to be carried out at the Kew Observatory, under NPL control, for the first decade before moving to the Teddington site in south-west London. A Morrison pendulum clock was installed at NPL as the reference standard for this work, adjusted to Greenwich Mean Time (GMT; mean solar time referenced to the Greenwich meridian) using time signals received over a telephone line from the Royal Observatory, Greenwich (ROG) [1].

The First World War generated an increased interest in standards for radio frequencies. During the 1920s David W. Dye developed a wavemeter based on an electrically-driven tuning fork that served as the primary frequency standard until 1934, when it was replaced by a quartz ring resonator that was also developed by Dye. Louis Essen joined NPL in 1929, working with Dye until the latter's untimely death in 1932. Essen was able to make a number of significant improvements to the Dye ring resonator, and an Essen ring resonator operating at 100 kHz became the primary standard in 1936.

During this period the ROG remained responsible for time in the UK. After 1927 NPL compared radio time signals from the ROG (via a transmitter at Rugby), Paris (transmitted from the Eiffel Tower) and Hamburg (transmitted from Nauen) with the ROG's time signal received over a telephone line. A temperature-controlled Shortt free-pendulum clock acted as the time standard at NPL, synchronised to the ROG signal, and remained in use until 1955 alongside the Essen ring resonator frequency standards. In 1939 an additional Essen ring resonator was made for the ROG, and became the first quartz clock installed at Greenwich. During the 1950s several Essen ring standards constructed to an improved design were sold to other laboratories, including the United States Naval Observatory (USNO), and an example is shown in figure 1.

* Peter Whibberley, Time and Frequency group, National Physical Laboratory
e-mail: peter.whibberley@npl.co.uk

The first caesium atomic clock

The concept of the atomic clock – that an atomic transition could serve as a natural standard of frequency – had been understood since the 1870s [2], but it would be several decades before the technology needed to construct a practical atomic clock became available. Much of the early work was carried out in the USA, in the molecular beam laboratory set up by Isidor I. Rabi at Columbia University, New York City. Working with a group of talented colleagues, Rabi developed a magnetic resonance method that could be used to observe transition frequencies in atomic and molecular beams. Rabi was aware that the method could form the basis of a frequency standard and discussed the idea with scientists at the National Bureau of Standards (NBS) in 1939, but the Second World War intervened and he first lectured in public on the possibility of atomic clocks in January 1945 [3].

The first true atomic clock was operated at NBS in August 1948, but it did not utilise Rabi's molecular beam technique. Harold Lyons, then in charge of the microwave standards group, instead constructed a clock based on a microwave absorption line in ammonia molecules contained within a spiral absorption cell. The performance of the ammonia clock, although comparable to the best quartz oscillators, was poorer than expected, and Lyons turned instead to the construction of a caesium beam clock. By 1952, the team led by Lyons and Jesse Sherwood were able to observe the caesium transition using the single-cavity Rabi resonance technique, but further development was needed in several key areas to produce an

operational clock. Staff changes, serious damage to the experiment and relocation to the new NBS site in Boulder, Colorado then combined to stall further progress.

At NPL, there had been considerable improvement of microwave techniques during the Second World War, one example being the development of resonant cavity wavemeters that could be used to stabilise microwave oscillators. Louis Essen was able to use a cavity wavemeter to obtain an improved estimate of the speed of light, and the expertise gained in high spectral purity frequency multiplication later proved to be of great benefit to atomic clock development. The post-war development of atomic clocks in the USA was viewed with much interest in the UK, not least by Harold Spencer Jones, the Astronomer Royal and Director of the Royal Greenwich Observatory (RGO, as it became known in 1948 prior to its move away from Greenwich). However, the RGO did not have the resources to build an atomic clock. At NPL, Louis Essen applied for permission to construct an atomic beam frequency standard following a visit to Lyons' laboratory at NBS in 1949, but the priority at that time was to complete a prototype stored-program computer known as Pilot ACE and the skilled staff needed for an atomic clock could not be spared.

Despite this setback, the Director of NPL, Sir Edward Bullard, was enthusiastic about the potential of the atomic clock and resources were gradually made available. In January 1953, Essen received funding from an aid programme administered by the US State Department to tour a number of laboratories in the USA involved in atomic clock development, including NBS and the laboratory of Professor Jerrold Zacharias at the Massachusetts Institute of Technology (MIT). On his return, Essen was granted approval to start work on an atomic clock. Although both the NBS and MIT caesium clocks were expected to be operational in the near future, there were good scientific reasons to have other independently-designed atomic clocks available for comparison, and one would be needed to serve as the new primary frequency standard for the UK. John ('Jack') Parry, an experienced experimental physicist, was assigned to work with Essen, and progress quickened.

The caesium atomic beam standard they constructed was based on an evacuated copper waveguide around 150 cm in length with a cross-section of 10 cm by 5 cm, as shown in figure 2. A significant improvement over the prototype NBS beam clock (which was also being incorporated into the NBS clock at around the same time) was the inclusion of an elongated U-shaped microwave cavity to probe the atomic beam in two short, well-separated interactions close to the ends

Figure 1:
An Essen ring
quartz resonator
standard.



of the cavity. The effect was to reduce the resonance line-width by increasing the total interaction time between the atoms and the microwave fields. This “separated oscillatory fields” method had been developed in 1949 by Norman Ramsey, a professor at Harvard University, and greatly increased the potential performance of the clock. The microwave synthesis chain developed by Essen and Parry made use of a 5 MHz quartz oscillator with very low phase noise that was obtained from the Marconi company.

By May 1955 the clock was complete (figure 3), and the broad tuning range of the microwave frequency enabled the atomic resonance line to be observed without too much difficulty (figure 4). Essen wrote later, “We invited the Director to come and witness the death of the astronomical second and the birth of atomic time”.

Determining the caesium frequency

Over the next few weeks, Essen and Parry evaluated the systematic biases of the clock frequency caused by environmental conditions such as temperature, pressure and background magnetic field strength. Each week the quartz oscillator standards were adjusted to agree with the unperturbed caesium frequency, then the ambient conditions were varied and the effect on the caesium clock frequency measured against the quartz standards. The tests showed that as expected the atomic clock was a much more stable timekeeper than the quartz clocks, except over short measurement periods, and more importantly it was also much more stable than astronomical time.

The next question to be addressed by Essen and Parry was the value of the caesium transition frequency. This value had to be measured empirically against the astronomical second, but which one? The time scales maintained by national observatories were representations of Universal Time (UT; the name chosen to replace GMT in 1928), based on timed observations of the Earth's angular position in space. Fluctuations in the Earth's rate of rotation had long been suspected, and the existence of seasonal variations was confirmed during the 1930s. Other apparent errors in observations were found to be caused by movement of the axis of rotation, or “polar motion”. From 1948 the RGO applied predicted corrections for polar motion to the observed Universal Time (called UT0) to produce a time scale that was more closely related to astronomical observations at any point on the Earth's surface (UT1). Soon the seasonal oscillations were also being corrected to give a smoother time scale (UT2), but even this was known to be affected by long-term changes in the Earth's rotation rate.

To satisfy the need for a more uniform time scale, an astronomical conference in 1950 recommended that a new time scale should be intro-

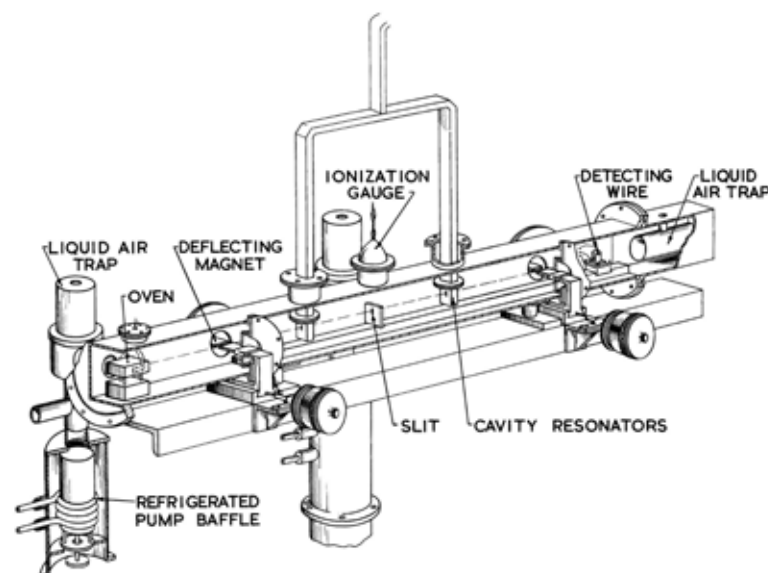
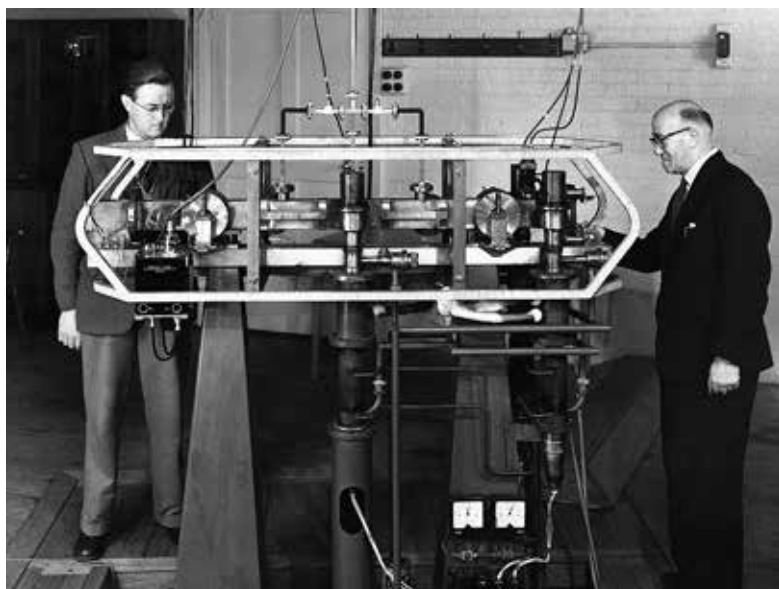


Figure 2:
Schematic drawing
of the internal
components of the
first NPL caesium
standard.

duced, to be called Ephemeris Time (ET) as it was based on mathematical models of the motions of astronomical bodies, and defined in terms of the Earth's orbit around the Sun. Since these models were based on astronomical observations extending from the mid-18th century to the beginning of the 20th century, the duration of the ephemeris second corresponded approximately to that of the UT second as it was around 1820. The International Astronomical Union (IAU) adopted ET in 1952, aiming to ensure by adopting a uniform time standard that timekeeping should continue to be defined by an astronomical phenomenon even after the introduction of atomic clocks [4]. In 1956 the International Committee for Weights and Measures (CIPM) selected the second of ET to be the base unit of time in the new International System of Units, and the decision was ratified in 1960 by the 11th General Conference on Weights and Measures (CGPM) as the first SI definition of the second.

Figure 3:
The first caesium
atomic frequency
standard at NPL,
with Louis Essen
(right) and Jack
Parry.



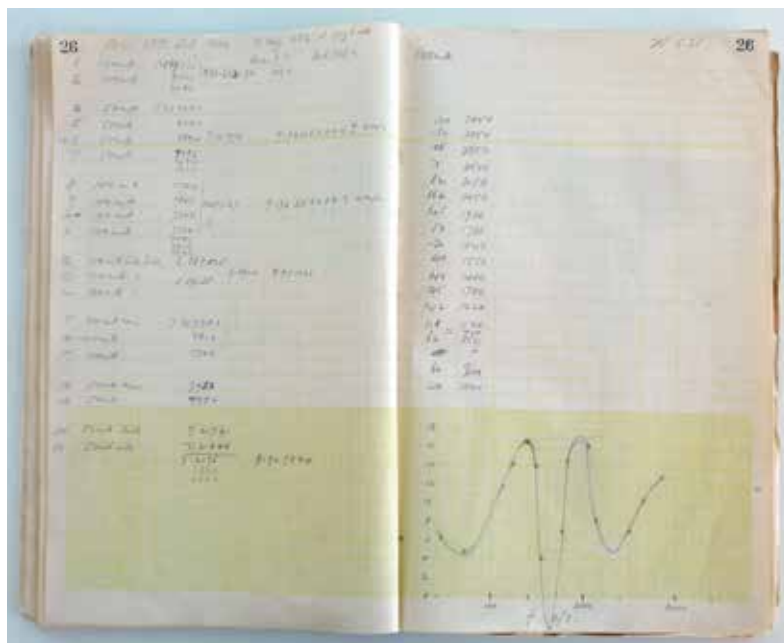


Figure 4:
Pages from Louis Essen's lab notebook dated 25 May 1955, showing the first observation of the caesium resonance line.

Despite these developments, the only astronomical time scale available at NPL in 1955 was the smoothed mean solar time (UT2) disseminated by the RGO. As soon as the caesium clock was operating reliably, Essen and Parry determined the frequency of the resonance against UT2, obtaining a value of $9192631830 \text{ Hz} \pm 10 \text{ Hz}$. They then made use of this number to adjust the Essen ring resonator standards each week so that they remained aligned in frequency with the caesium clock. This was the first occasion when a clock surpassed the performance of astronomical time, referenced to a quantum transition frequency that is believed to be an unchanging natural constant. It also marked the start of the first atomic time scale.

Figure 5:
A USNO moon camera, of the type developed and used by William Markowitz to realise Ephemeris Time.

Aware of the considerable interest in atomic clock developments, Essen and Parry quickly prepared a short report on the first operation of the NPL clock, including the result of the frequency

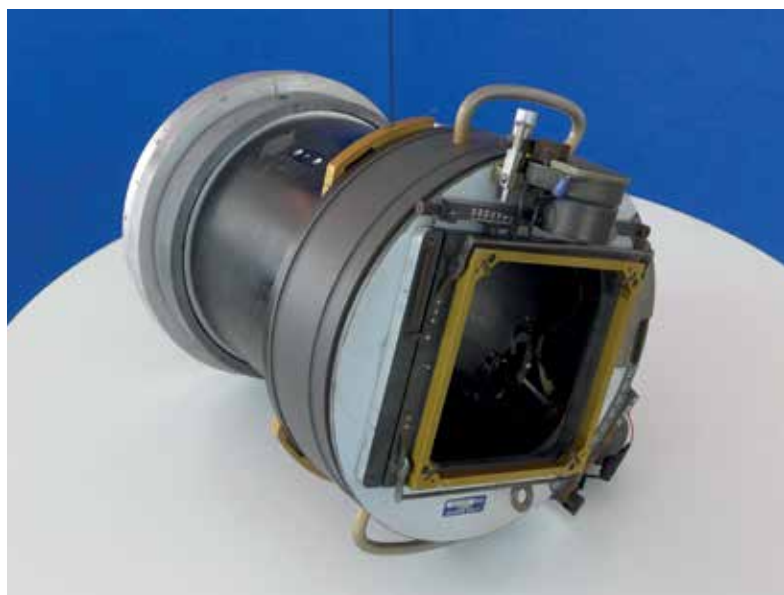
measurement, which was published in *Nature* on 13 August 1955 [5]. The claimed uncertainty of 1 part in 10^9 was already limited by the realisation of Universal Time, but the report pointed out that improvements to the clock electronics would increase its precision considerably. By the following summer the uncertainty of the clock had been improved. In a second paper that also provided a detailed description of the clock [6], the caesium frequency was reported to be $9192631845 \text{ Hz} \pm 2 \text{ Hz}$ compared to the RGO's realisation of UT2 averaged over the period from June 1955 to June 1956.

A determination of the caesium transition frequency compared to the new ET definition of the second would be considerably more difficult to achieve. In August 1955, Essen attended a meeting of the IAU General Assembly, which confirmed the adoption of Ephemeris Time despite hearing Essen's report of the initial results from the NPL atomic clock. The Director of the Time Service Department at the USNO, William Markowitz, was also present. Markowitz had developed a novel type of Moon camera specifically to determine Ephemeris Time (figure 5), and he and Essen agreed to collaborate on a measurement of the caesium frequency in terms of ET. Over the next 3 years, Markowitz carried out an extended series of observations of the Moon's position using the camera, and by fitting the results to a mathematical model of the Moon's motion he was able to obtain a precise realisation of the ephemeris second. The astronomical measurements were compared regularly with the frequency of the NPL caesium clock using long-range radio time signals that could be received in both locations.

The final value obtained for the frequency of the caesium transition was $9192631770 \text{ Hz} \pm 20 \text{ Hz}$ [7]. The ephemeris second was therefore slightly shorter than the UT2 second at that time. Although other atomic clocks had been brought into operation during the comparison period, there would have been no benefit from including them in the analysis as the result was dominated by the uncertainties of the astronomical observations and the radio time signal comparisons. The measurement campaign has never been repeated, and the result obtained by Markowitz and Essen was adopted for the redefinition of the SI second in 1967.

Other developments at NPL leading towards the 1967 redefinition

The first caesium beam standard at NPL was essentially a prototype, and it included a number of features intended to simplify bringing the clock into operation and to investigate potential systematic biases. By late 1955, Essen and Parry were



constructing a self-contained, transportable experimental clock that was mounted on a trolley. Its vacuum chamber was constructed of steel to assess the effects of this magnetic material on the clock's performance. In 1958, the transportable clock was demonstrated at the "Pendulum to Atom" public exhibition in London on the development of precise timekeeping (figure 6) [8].

The experience gained from both the first clock and the transportable clock was incorporated in the design of a new caesium beam standard, known as NPL2. This clock used a much longer beam tube than the earlier clocks, around 5 m in length, and it was orientated vertically with the oven at the bottom (figure 7). This arrangement reduced the alignment complications caused by the parabolic path of the atoms, and in principle gave rise to a slight narrowing of the linewidth due to the gravitational slowing of the beam. Another novel feature of NPL2 was the inclusion of microwave cavities for both rubidium and caesium, to allow the resonant frequencies of the two elements to be compared directly. The new clock was operational by the summer of 1958. It was found to suffer from a frequency offset that was traced to cavity asymmetry, and modifications to investigate the effect included the installation of ovens and detectors at both ends (for caesium only) to allow bi-directional operation. NPL2 was upgraded several times to improve its performance, and continued in use as the primary frequency standard at NPL until the early 1970s [9].



Figure 6: The transportable caesium beam clock constructed at NPL in 1955–56, shown on display to the public in London in 1958 as part of an exhibition on the development of precise timekeeping called "Pendulum to Atom".

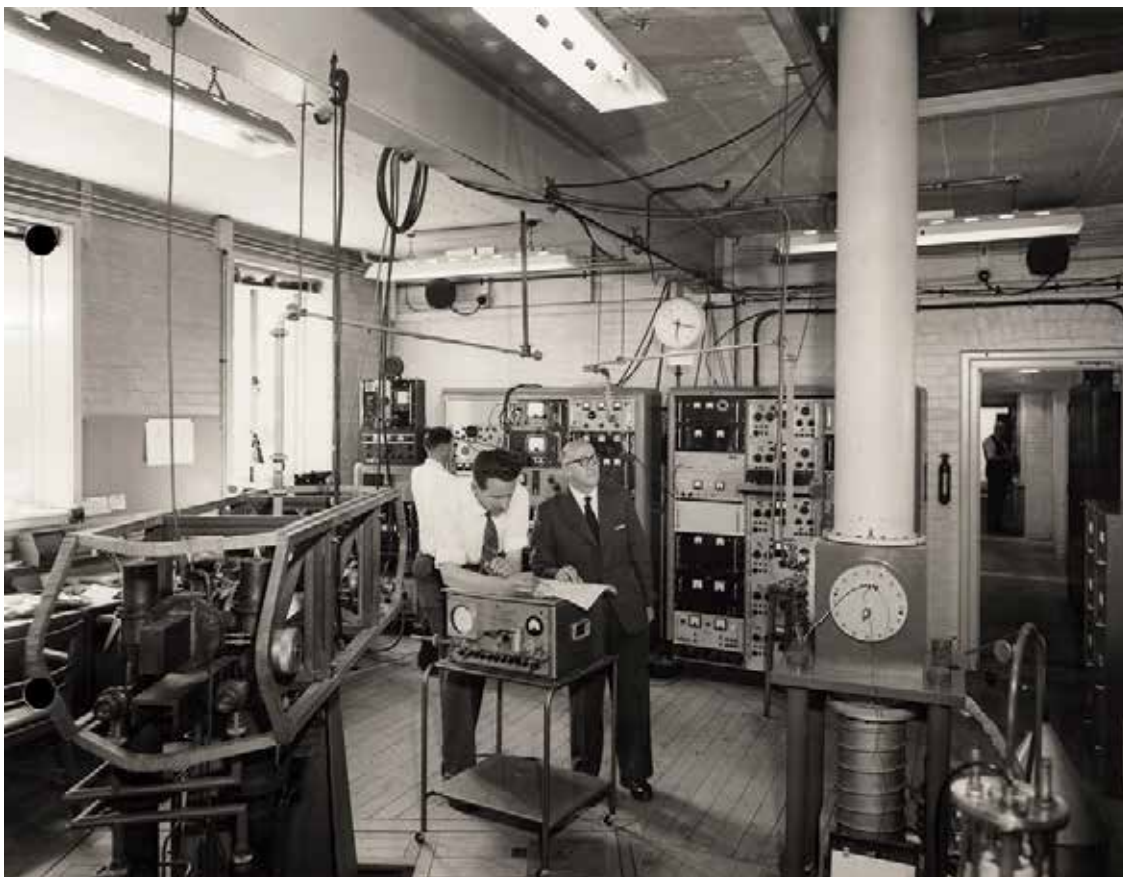


Figure 7: The NPL atomic clock laboratory in 1962, showing the first clock on the left and the vertical-beam clock NPL2 on the right. Denis Sutcliffe and Louis Essen are standing in the centre.

Several conditions had to be met before the idea of redefining the second in terms of an atomic transition would gain widespread acceptance [10]. One was the need to demonstrate that caesium beam standards designed and constructed independently in different laboratories did indeed agree in frequency to within their estimated uncertainties. The first successful atomic clock in the USA was the “Atomichron”, manufactured commercially by the National Company in Massachusetts to the design of Zacharias and available commercially from late 1956. Initial measurements between a pair of Atomichrons and the NPL clock by the common-view of radio signals indicated a significant frequency difference between the two designs. Two Atomichrons were then shipped to the NPL and operated alongside the NPL clock so that a direct comparison could be made, this time resulting in much closer agreement, and the cause of the discrepancy in the original measurements was traced to the link analysis rather than the clocks [11].

As more caesium beam standards were brought into operation in the USA and other countries, long-distance frequency comparisons by observing low frequency (LF) or very low frequency (VLF) radio signals controlled by stable oscillators became routine. In 1950 NPL became responsible for the MSF transmissions from Rugby Radio Station in central England on 2.5, 5 and 10 MHz, and on 60 kHz for 1 hour each day initially (figure 8). The 60 kHz signal proved to be well-suited to transatlantic comparisons, and from June 1955 onwards the radiated frequency, controlled by a quartz oscillator, was maintained in agreement with that of the NPL caesium clock

to improve its stability. This was the first time that a standard-frequency transmission had been based on an atomic clock.

A key development took place in 1959 with the introduction of coordination between the time signal transmissions in the UK and the USA. The signals were linked to the atomic time scales at the NPL and the NBS respectively, and agreement was reached to apply the same constant frequency offset, changed each year, to both signals to maintain the time markers in close alignment with UT2. When required, additional adjustments by small time steps were applied on agreed dates. The common broadcast time scale soon became known as Coordinated Universal Time (UTC), and other countries operating radio time signals adopted the same procedure so that their transmissions also became representations of UTC. During the early 1960s, the various UTC transmissions were typically synchronised at the millisecond level, and were maintained to within about 20 ms of UT2.

The availability of long-range radio signals also allowed integrated atomic time scales to be formed by averaging clocks in different laboratories. An important step was taken in 1958 when Markowitz at the USNO introduced the atomic time scale A.1, adjusted using signals derived from the NPL caesium clock and based on the recently-determined value for the caesium transition frequency in terms of Ephemeris Time. The scale was defined so that the A.1 time agreed with the UT2 time at 0h on 1 January 1958. Soon the caesium clocks being brought into operation in other institutes were incorporated into A.1, and by 1962 there were 9 laboratories participating in the system, including those in Canada, France and Switzerland as well as at the NPL [12]. Later, the A.1 scale was modified to become an internal USNO time scale.

The Bureau International de l’Heure (BIH), located at the Paris Observatory, had been established in 1912 to realise a definitive Universal Time based on measurements from observatories around the world [13]. It took an early interest in atomic time, and in July 1955 it started to construct a time scale named AT, referenced to the NPL caesium standard by observations of the MSF transmission frequency. From July 1956, clocks in other laboratories were included in a mean atomic time scale called AM, and in 1963 a second scale A3 was formed using only the primary frequency standards constructed at metrology institutes. At first only NPL, NBS and LSRH (Laboratoire Suisse de Recherches Horlogères) operated such primary standards, but within a few years other laboratories began to contribute to A3. The methods employed by the BIH for averaging the independent atomic time scales were steadily improved,

Figure 8:
One of the 250 m masts at Rugby radio station, one of several in Europe and the USA that transmitted the long-range standard-frequency and time signals used to compare atomic clock frequencies and national time scales.



leading to the renaming of the BIH time scale as International Atomic Time (TAI) in 1971 [14].

One further example will serve to illustrate the range of work being carried out in time and frequency metrology at NPL during the early 1960s. The accuracy of long-distance time comparisons using low-frequency radio signals was limited to around 1 ms, and more precise methods were needed to improve the coordination of national time scales. The launch of the first geostationary communication satellite, Telstar I, in July 1962 provided the opportunity to exchange pulsed timing signals simultaneously in both directions between the UK and the USA. James McA. ('Mac') Steele at the NPL collaborated with William Markowitz and his colleagues at the USNO in an experiment that compared the RGO and USNO time scales, with impressive results. The satellite link was found to be accurate to around 1 microsecond, with the potential for further improvement, and the trial has led to the routine use of two-way satellite time and frequency transfer (TWSTFT) between national timing centres today [15].

Summary

Despite its relatively limited resources, NPL made considerable contributions to the development of atomic timekeeping, and to the adoption of the caesium transition frequency as the definition of the SI second. In 1955, Louis Essen and Jack Parry succeeded in constructing and operating a caesium atomic beam frequency standard, which was the first clock with substantially better stability than the rotation of the Earth and also the first true quantum standard. In a collaboration with the US Naval Observatory, the NPL clock was employed to determine the caesium transition frequency compared to Ephemeris Time, and in 1967 the value obtained was adopted as the new definition of the SI second.

Other, less well-known, activities in time and frequency at NPL during this period included the construction of a transportable caesium beam clock and a vertical caesium beam primary frequency standard. These clocks were run regularly to adjust the quartz resonator generating the NPL time scale, the first continuously available atomic time scale. NPL also operated stable low-frequency transmissions used for long-distance frequency comparisons between atomic clocks, and investigated novel comparison techniques such as the first two-way time transfer using a geostationary communication satellite.

References

- [1] E. Pyatt, "*The National Physical Laboratory, A History*", Adam Hilger, Bristol, 1983
- [2] W. Snyder, "*Lord Kelvin on Atoms as Fundamental Natural Standards (for Base Units)*", IEEE Transactions on Instrumentation and Measurement 22(1), **99**, 1973
- [3] R. Essen, "*The Birth of Atomic Time*", Fastprint Publishing, Peterborough, 2015
- [4] G. A. Wilkins, "*A personal history of the Royal Greenwich Observatory at Herstmonceux Castle 1948–1990*", Sidford, Devon 2009
- [5] L. Essen and J. V. L. Parry, "*An Atomic Standard of Frequency and Time Interval: A Caesium Resonator*", Nature 176, 280–282, 13 August 1955
- [6] L. Essen and J. V. L. Parry, "*The Caesium Resonator as a Standard of Frequency and Time*", Philosophical Transactions of the Royal Society of London **A 250**, 45–69, 1957
- [7] W. Markowitz, R. G. Hall, L. Essen and J. V. L. Parry, "*Frequency of caesium in terms of ephemeris time*", Phys. Rev. Lett. **1**, 105, 1958
- [8] R. Essen, "*Two clocks that changed the world, the birth of atomic timekeeping*", Antiquarian Horology **34**(2), 219–234, 2013
- [9] R. Essen, "*The First Long-Beam Atomic Clock*", The Horological Journal, 168–172, April 2017
- [10] S. Leschiutta, "*The definition of the 'atomic' second*", Metrologia **42**, S10–S19, 2005
- [11] J. A. Pierce, "*Recent long-distance frequency comparisons*", IRE Transactions on Instrumentation, 207–210, 1958
- [12] W. Markowitz, "*The atomic time scale*", IRE Transactions on Instrumentation, 239–242, 1962
- [13] B. Guinot, "*History of the Bureau International de l'Heure*", Astronomical Society of the Pacific Conference Series 208, 2000
- [14] B. Guinot and E. F. Arias, "*Atomic time-keeping from 1955 to the present*", Metrologia **42**, S20–S30, 2005
- [15] J. McA. Steele, W. Markowitz and C. A. Lidback, "*Telstar Time Synchronization*", IEEE Transactions on Instrumentation and Measurement **13**(4), 164–170, 1964

A Historical Review of U. S. Contributions to the Atomic Definition of the SI Second

Michael A. Lombardi*

1. Introduction

The second, the base unit of time interval, is one of the seven base units in the International System (SI). Because time interval and its reciprocal, frequency, can be measured with more resolution and less uncertainty than any other physical quantity, the second holds a place of preeminence in the world of metrology. The precise realization of the SI second that we benefit from today was made possible by the development of atomic frequency standards and clocks; devices that fundamentally changed the way that time is measured and kept. Before atomic clocks, the second was defined by dividing astronomical events, such as the solar day or the tropical year, into smaller parts. This permanently changed in 1967, when the SI second was redefined as the duration of 9 192 631 770 periods of the electromagnetic radiation that causes ground state transitions in the cesium atom [1]. The new definition meant that seconds were now measured by counting oscillations of electric fields that cause atoms to change state, and minutes and hours were now multiples of the second rather than divisions of the day.

The benefits of atomic timekeeping to modern society have been immense and are difficult to overestimate. Many technologies that we now take for granted, such as global navigation satellite systems, mobile telephones, and the “smart grids” that provide our electric power, depend upon atomic clock accuracy. This makes it easy to forget that the era of atomic timekeeping began less than an average lifetime ago. The current year (2017) marks a half-century since the International System of Units (SI) adopted a definition of the second based on an atomic transition [1]. To commemorate this 50th anniversary, this paper provides a historical review of work that contributed to the atomic definition of the SI second in 1967. Its primary focus is on work conducted in the United States of America.

2. The First Suggestion of Atomic Clocks and Early Research

The possibility of an atomic clock was first suggested in Europe in the 19th century. James Clerk Maxwell, a Scottish physicist, was likely the first person to recognize that atoms could be used to keep time. In an era during which the Earth’s rotation was the world’s timekeeping reference, Maxwell remarkably suggested to Peter Guthrie Tait (a childhood friend) and William Thomson (later to be known as Lord Kelvin) that the “period of vibration of a piece of quartz crystal of specified shape and size and at a stated temperature” would be a better absolute standard of time than the mean solar second, but it would still depend “essentially on one particular piece of matter, and is therefore liable to all the accidents, etc. which affect so-call National Standards however carefully they may be preserved, as well as the almost insuperable practical difficulties which are experience when we attempt to make exact copies of them.” [2] Atoms would work even better as a natural standard of time. As Thomson and Tait wrote in their *Elements of Natural Philosophy*, first published in 1879:

“The recent discoveries ... indicate to us natural standard pieces of matter such as atoms of hydrogen or sodium, ready made in infinite numbers, all absolutely alike in every physical property. The time of vibration of a sodium particle corresponding to any one of its modes of vibration is known to be absolutely independent of its position in the universe, and will probably remain the same so long as the particle itself exists.” [2, 3]

It was to be some six decades after these prophetic mentions of atomic clocks before the first experiments were performed. The experiments were finally made possible by the rapid advances in quantum mechanics and microwave electronics that took place before, during, and after World War II. Most of the fundamental concepts that would eventually lead to an atomic

* Michael Lombardi, Time and Frequency Division, National Institute of Standards and Technology, Boulder, Colorado, USA, lombardi@nist.gov

clock were developed by Isidor Isaac Rabi and his colleagues at Columbia University in New York City, beginning in the 1930s [4, 5]. Born in 1898 in Rymanów, Galicia, then part of Austria-Hungary (now Poland), Rabi emigrated to the United States as an infant, grew up in New York City, and went on to attend Columbia University. After initially studying electrical engineering and chemistry, he eventually received a doctorate in physics. In 1929, he joined the Columbia faculty. Shortly afterward, he built a molecular beam apparatus, and his new molecular beam laboratory soon began to attract other talented young physicists who would go on to play key roles in atomic clock history, including Sidney Millman, Jerrold Zacharias, Polykarp Kusch, and Norman Ramsey.

Rabi invented the atomic and molecular beam resonance method of observing atomic spectra in 1937, groundbreaking work for which he would receive the Nobel Prize in physics in 1944. As early as 1939, he had informally discussed applying his molecular beam magnetic resonance technique as a time standard with scientists at the National Bureau of Standards (NBS), then located in Washington, D.C. However, World War II soon followed, and Rabi, like most scientists of that era,

halted his own research to assist in the war effort. During the war, he led a group of scientists at the Massachusetts Institute of Technology (MIT) in Cambridge, Massachusetts, who helped in the development of radar. During the final year of the war, Rabi publicly discussed the possibility of atomic clocks for the first time in a lecture given to the American Physical Society and the American Association of Physics Teachers on January 20, 1945 [6]. The following day, the *New York Times* covered the story (Fig. 1) with a lead sentence referring to the “most accurate clock in the universe.” [7]

The *New York Times* story referred to a “cosmic pendulum,” a metaphorical reference acknowledging that the basic principles of atomic clocks were straightforward and were already well understood by the early researchers. In short, because all atoms of a specific element are identical (ignoring isotopic differences), they should undergo a transition in their energy level when stimulated with electromagnetic radiation at a specific frequency. This meant that an atom could potentially serve as a nearly perfect resonator for a clock. The frequency that caused the energy transition could be used as a standard of frequency, and the periods of this frequency could be counted to define a standard for time interval. It easily followed that a clock referenced to an atomic resonator would be far more accurate than previous clocks that were based on mechanical oscillations of a pendulum or a quartz crystal.

Although the concept of an atomic clock was seemingly simple, its implementation was exceedingly difficult, particularly with the technology that existed in the 1940s. Atomic transitions occurred in the microwave region, at much higher frequencies than had previously been measured, and systems and electronics were needed that could measure these microwave frequencies. Determining the resonance frequency of a specific element was necessary so that it could be related to the second, the base unit of time interval. The next step would be to use the atomic resonance frequency to control the frequency of a physical oscillator, such as a quartz crystal oscillator, by locking it to atomic resonance. Finally, to make the output of an atomic clock usable, the frequency of the locked quartz oscillator would need to be divided into frequencies low enough to use for timekeeping.

The first trick was to be able to measure the resonance frequency of an atom, which is derived from its quantized energy levels. The laws of quantum mechanics dictate that the energies of a bound system, such as an atom, have certain discrete values. An electromagnetic field generated at a specific frequency can boost an atom from one energy level to a higher one, or an atom at

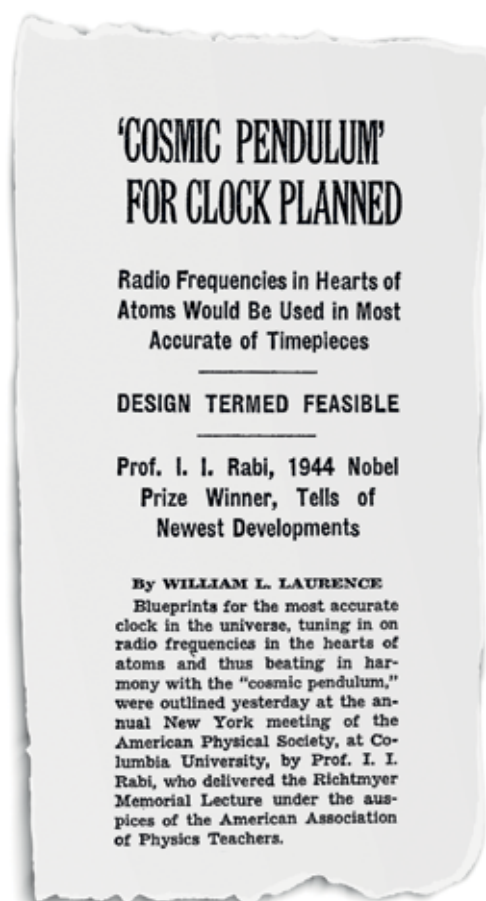


Figure 1L: The first public suggestion of an atomic clock, *New York Times*, January 21, 1945, p. 34 (first paragraph of article).

a high-energy level can drop to a lower level by emitting energy. The resonance frequency (f_o) of an atomic oscillator is the difference between the two energy levels, E_1 and E_2 , divided by Planck's constant, h :

$$f_o = \frac{E_2 - E_1}{h}. \quad (1)$$

The energy absorbed or emitted by the atoms is spread across a small frequency range that surrounds f_o , and therefore it does not exist at f_o alone. This spread of frequencies, Δf_a , is known as the resonance width, or linewidth. The ratio of the resonance frequency to the resonance width is known as the quality factor, Q , where

$$Q = \frac{f_o}{\Delta f_a}. \quad (2)$$

Measuring the atomic linewidth was another challenging task faced by early researchers. However, because their resonance frequencies were so high, it was obvious that an atomic clock, even with a broad linewidth, would have a much higher Q than any previous clock and would therefore have a potentially huge advantage in both accuracy and stability.

It seems likely that Rabi and his colleagues expected the energy transition in cesium (^{133}Cs) atoms to be the resonator for the first atomic clock [8]. Millman and Kusch (who, like Rabi, would become a Nobel Laureate in 1955), had first measured the cesium resonance frequency at Columbia in 1940, estimating the frequency of the hyperfine transition as 9191.4 megacycles. This was relatively close to the number that would later define the second [9]. Fate would intervene, however, and the first atomic clock was not based on cesium, but instead was an ammonia absorption device built at NBS.

3. The First Atomic Clock, 1948

Harold Lyons was certainly not the first scientist to perform spectroscopic frequency stabilization experiments or to engage in atomic clock research, but he and his colleagues at NBS were the first to build a working atomic clock. Born in Buffalo, New York, in 1913, Lyons was a graduate student at the University of Michigan (he obtained a Ph.D. in nuclear physics in 1939) when Rabi and his colleagues at Columbia first discussed employing their molecular beam resonance technique as a time standard. Lyons joined NBS in 1941 and became head of the microwave standards group in 1948. There, he and his colleagues moved rapidly past the pure experimental stage and began constructing a working atomic clock. After a period

of intense research and development beginning in the spring of 1948, the new clock (Fig. 2) was first operated on August 12, 1948, and was publicly demonstrated in January 1949 [10].

The construction of an atomic clock in 1948 was a significant scientific and engineering accomplishment that required many obstacles to be overcome. Because he had chosen the absorption line of ammonia gas as the source of frequency, Lyons's first obstacle was measuring the inversion transition frequency, or what Lyons called, in keeping with Thomson and Tait [2], the "vibration" [11], of the atoms inside ammonia molecules. The ammonia gas was stored in an absorption cell, which was designed as a 30 ft long (~9 m long) copper tube that was wrapped into a spiral. The measurement involved multiplying a 100 kHz signal from quartz oscillator several times. The first stage consisted of a frequency multiplication chain built from low-frequency tubes that resulted in a frequency of about 270 MHz. The second stage consisted of a frequency multiplying klystron (a linear beam vacuum tube), wherein frequency was modulated by a 13.8 MHz oscillator that multiplied the signal to about 2.983 GHz. Finally, the signal was multiplied to approximately 23.87 GHz by use of a silicon crystal rectifier. This frequency was fed to the ammonia absorption cell and tuned until

Figure 2: Dr. Harold Lyons (right), inventor of the ammonia absorption cell atomic clock, observes, while Dr. Edward U. Condon, the director of the National Bureau of Standards, examines a model of the ammonia molecule (1949). The clock atop the equipment racks was driven by a 50 Hz signal obtained by dividing the signal from a quartz oscillator that was locked to the ammonia absorption frequency. However, its primary function was to let the world know that the invention was indeed a clock.



the signal reaching a silicon crystal detector at the end of the cell showed a spectral line (in the form of a dip, or inverted peak) when displayed on an oscilloscope (Fig. 3). The line indicated that the incoming frequency from the multiplication chain

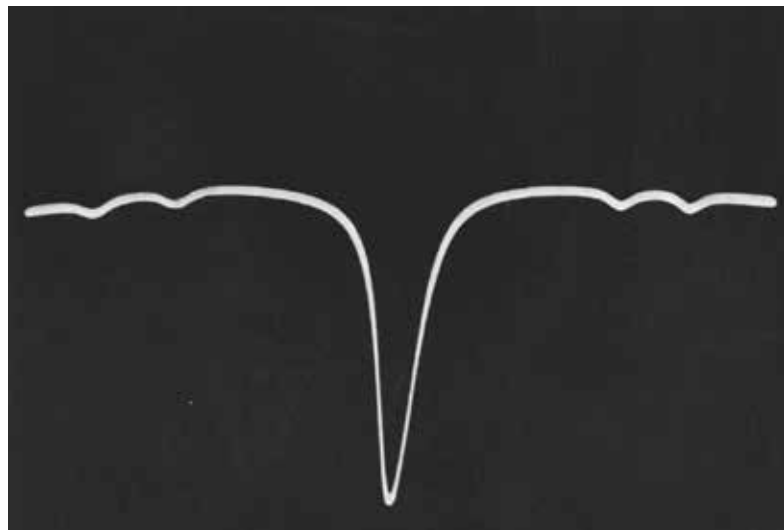


Figure 3:
The spectral line (shown on an oscilloscope) that occurred when the incoming microwave signal from the frequency multiplication chain matched the absorption frequency of the ammonia molecule.

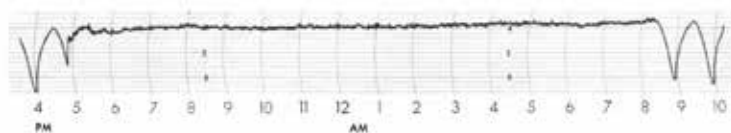


Figure 4:
A strip chart showing a period (from 5 p.m. to 8 a.m.) where the quartz oscillator was locked to the absorption frequency of ammonia.

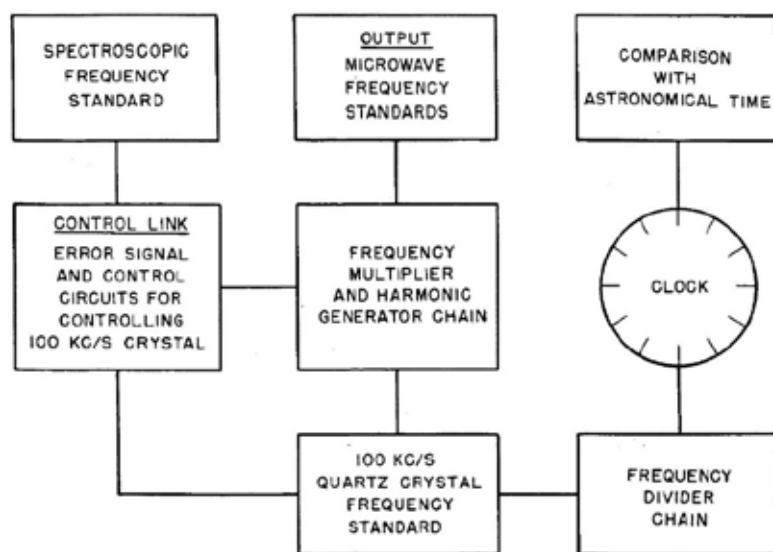


Figure 5:
A simplified block diagram of the first atomic clock.

now matched the ammonia absorption frequency. When the line occurred, an electrical pulse was generated [11–14].

The next obstacle faced by Lyons and his group was to develop a method to automatically adjust the quartz oscillator frequency until it matched the frequency of the ammonia absorption line. This was done by taking a 12.5 MHz frequency from the quartz oscillator frequency multiplication chain and mixing it with the signal from the 13.8 MHz frequency-modulated oscillator. When the frequency from the mixer matched its expected value, another electrical pulse was generated. The time interval between the two generated pulses (the pulse from the absorption cell and the pulse from the mixer) was continuously measured. If the time interval was increasing, it indicated that the quartz oscillator frequency was increasing, and vice versa. The time interval measurements were used to generate a control signal that adjusted the quartz oscillator whenever necessary to keep it locked to the ammonia absorption frequency [11, 13].

Figure 4 shows a strip chart where the quartz oscillator was locked to the ammonia absorption frequency from 5 p.m. to 8 a.m. Before and after this period, the quartz oscillator was “unlocked” (free running). Each small division on the strip chart represents a frequency variation of less than 1 part per 10 million ($1 \cdot 10^{-7}$).

A third obstacle involved dividing the locked quartz oscillator frequency into low frequencies that could be used for timekeeping [15]. The first atomic clock included dividers that produced two output signals, a 50 Hz signal that drove an ordinary synchronous motor clock, and a 1000 Hz signal that was used to drive a special synchronous motor clock that was designed for comparisons with astronomical time (then the world standard for timekeeping) to within 5 ms. Figure 5 is a simplified block diagram showing how the clock worked.

A final technical obstacle faced by Lyons was his quest to make his ammonia clock more accurate by increasing the quality factor, Q , as previously discussed and shown in Eq. (2). It was already known that oscillators with higher Q values would be potentially more stable and accurate, so it was desirable to increase Q as much as possible, either by using an atomic transition where the frequency was as high as possible, or by making the line-width as narrow as possible. In the case of the ammonia clock, the resonance frequency was high (23.8 GHz), but the resonance width, or the range of frequencies over which the oscillator could resonate, represented by the width of the line in the absorption frequency shown in Fig. 3, was not especially narrow. Lyons and his colleagues had little success when attempting to reduce the line-width. Thus, the actual Q factor of the ammonia

spectrum line was typically measured in a range from 50 000 to 500 000, approximating that of the best quartz crystal oscillators of that era, although it was much more constant and stable [12].

In a metrological sense, Lyons and his colleagues had advantages over university researchers, because NBS already had several decades of experience in maintaining, disseminating, and measuring standard frequency signals. This allowed Lyons to accurately measure his new clock against the national standard of frequency, the quartz oscillators operated by NBS that were periodically calibrated against the astronomical time standards maintained by the U.S. Naval Observatory. He quickly learned from these measurements that the potential accuracy of the ammonia clock was limited. Two versions of the ammonia clock were built, with the best reported accuracy about 2 ms/day ($2 \cdot 10^{-8}$). Work on a third version was halted when it became apparent that cesium beam techniques were likely to provide a significant increase in accuracy. Even though the cesium resonance frequency was much lower (~ 9.2 GHz), there were numerous options available for decreasing the linewidth, leading to Q factors that were much higher than those obtained with molecular absorption methods. Thus, by the early part of 1950, Lyons and his NBS colleagues had turned their attention away from ammonia clocks, and with some assistance from Kusch at Columbia University, had begun constructing a cesium beam clock [16].

4. Cesium Atomic Clocks

Cesium had several characteristics that made it suitable for use in an atomic clock. A silvery-white ductile metal, cesium is the softest element listed on Mohs hardness scale (0.2 MPa) and one of the few elemental metals that becomes a liquid at near room temperature (28.4°C); only mercury has a lower known melting point. Cesium atoms are also relatively heavy and thus move at a relatively slow speed of about 130 m/s at room temperature in an atomic beam. This allows cesium atoms to be observed for a longer period than hydrogen atoms, for example, which travel much faster, about 1600 m/s, at room temperature. Cesium also has a higher resonance frequency (~ 9.2 GHz) than other atoms that were later utilized in atomic clocks, such as rubidium (~ 6.8 GHz) and hydrogen (~ 1.4 GHz).

The first cesium clock built by Lyons and Jesse Sherwood at NBS utilized Rabi's magnetic resonance technique. Sherwood reported initial results at a 1952 meeting of the American Physical Society [17], but the resonance width of the clock was too wide (~ 30 kHz) to serve as an accurate time standard (the Q was only 300 000). Once cesium was selected as the atom of choice, the resonance

frequency could not be changed, so the focus turned to making the resonance width as narrow as possible. Narrowing the resonance width required a better way of interrogating the cesium atoms.

Rabi's original magnetic resonance technique interrogated the atoms with one microwave pulse. Methods were developed to increase the Q by applying longer microwave pulses, but these schemes did not work as expected due to subtle technical side effects; the long interaction period between the atoms and microwave field subjected the frequency of the clock to Doppler shifts and other uncertainties. A breakthrough occurred in 1949, when Norman Ramsey, a former student of Rabi who was then a professor at Harvard University, invented the separated oscillatory field method, an improvement that would be critically important to future atomic clock designs. Ramsey's new method interrogated the atoms with two short microwave pulses, separated by some distance in an atomic beam. Applying the oscillating field in two steps accomplished the goal of narrowing the resonance width. It also reduced the sensitivity to microwave power fluctuations and magnetic fields by factors of 10 to 100 or more, and it eliminated most of the Doppler effect. In short, Ramsey made it possible to build much more accurate atomic clocks [5, 18]. Some four decades later, in 1989, Ramsey received a Nobel Prize in physics for this work.

The NBS team quickly redesigned their clock to employ Ramsey's new technique of separated oscillating fields by applying microwave pulses in separated regions along an atomic beam. This reduced the resonance width to just 300 Hz, increasing the Q to about 30 million. Encouraged by these results, Lyons predicted that the clock would eventually be accurate to $1 \cdot 10^{-10}$, meaning that it would be able to keep time to within about 10 $\mu\text{s/day}$ [19]. His prediction eventually came true, but it happened much later than expected. In 1953, NBS interrupted their atomic clock program, a decision made partially for budgetary and political reasons and partially because the agency elected to focus on other areas. Sherwood had resigned in 1952 [10], and Lyons moved to the new NBS laboratories in Boulder, Colorado, in 1954. A year later, he left NBS entirely, accepting a position at Hughes Aircraft Company in Culver City, California.

The cesium clock was disassembled and moved from Washington, D.C., to Boulder, where it was reassembled by a new team of researchers that was led by Richard Mockler and included Roger Beehler and Chuck Snider. It eventually became a working time standard in 1958. Then known as NBS-1 (Fig. 6), the clock finally achieved Lyons' predicted accuracy [20, 21].

The length of the microwave interrogation cavity (commonly known as the Ramsey cavity) in

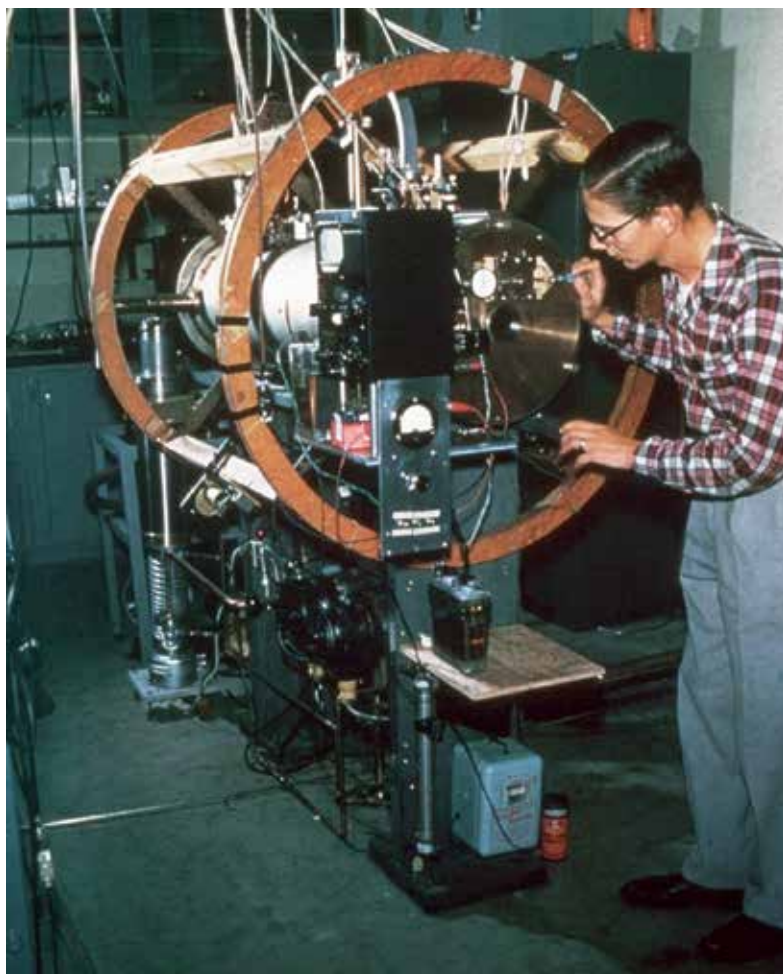
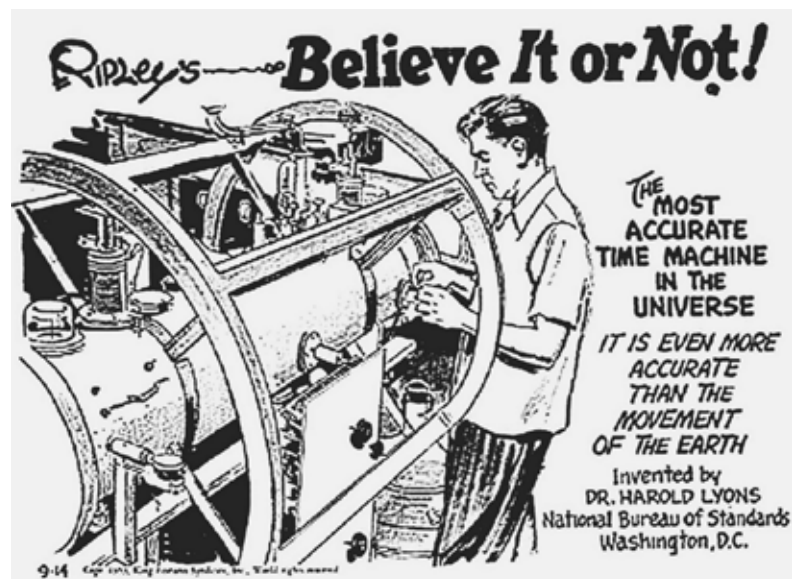


Figure 6:
Roger Beehler of
NBS with NBS-1.

Figure 7:
A 1953 cartoon
featuring the first
cesium clock devel-
oped at NBS.



Ironically and perhaps appropriately, *Ripley's Believe it or Not*, a popular syndicated cartoon featured in U.S. newspapers, credited Lyons with the invention of the atomic clock in 1953, but the drawing was not of the ammonia clock, but rather of the cesium beam clock that would later become NBS-1 (Fig. 7).

The disruption of the atomic clock program and a work stoppage of at least several years caused NBS to lose their large early lead in the cesium clock race and the National Physical Laboratory (NPL) in the United Kingdom became the clear winner. The NPL effort was led by Louis Essen, who had previously worked on quartz clocks and measurements of the speed of light before becoming interested in atomic timekeeping. Beginning in 1949, Essen made several trips to the United States, where he met with early atomic clock researchers, including Rabi and his colleagues at Columbia and Jerrold Zacharias of MIT. Essen was anxious to apply their work to timekeeping and appealed to Sir Edward Bullard, the director of NPL, to start a new program designed to build an atomic clock. After several delays, work on a cesium clock began at NPL in 1953, the same year when work was essentially stopped at NBS, and it progressed rapidly enough for a working time standard to be placed in operation in June 1955. The linewidth was 340 Hz, and the original reported accuracy was $1 \cdot 10^{-9}$ (100 $\mu\text{s/day}$) [25, 26]. The NPL cesium clock then famously played a key role in the experiment that led to the atomic definition of the SI second, as discussed in the next section.

5. Defining the Atomic Second by Comparison to Ephemeris Time

As noted in the introduction, prior to the invention of atomic clocks, the second was defined by dividing the period of an astronomical event into a shorter time interval. For example, the second was once defined by dividing the average period of one revolution of Earth on its axis. The mean solar second was equal to $1/86\,400$ of the mean solar day. To create a more stable unit of time interval, the second was redefined in 1956 as $1/31\,556\,925.9747$ of the tropical year 1900, a time interval known as the ephemeris second. The ephemeris second was indeed more stable than the mean solar second, but it was nearly impossible to use as a time reference in a laboratory, and was thus of little use to metrologists or engineers. In retrospect, it seems unfortunate that another astronomical definition of the second was accepted during a period when working atomic clocks were already being demonstrated [27, 28]. A clean transition from the mean solar second to the atomic second would certainly have made more sense. Doomed from the start, the ephemeris

second would be easy for historians to dismiss had it not served as the comparison reference for the definition of the atomic second.

Ephemeris time was determined by measuring the position of the Moon with respect to several surrounding stars. The most accurate Moon observations had been recorded at the U.S. Naval Observatory (USNO) in Washington, D.C., by the astronomer William Markowitz. Born in Austria in 1907, Markowitz's family emigrated to the United States when he was three years old. He obtained a Ph.D. in astronomy from the University of Chicago in 1931, and he worked at USNO from 1936 until 1966. By 1952, Markowitz was performing Moon observations with a sophisticated dual-rate Moon camera of his own design [28, 29].

Collecting Moon observations was a tedious task and results were obtained very slowly. Bullard, the NPL director, wrote in 1955 that it would take four years of Moon observations to determine time with the same accuracy as their new cesium clock. He also noted that "atomic clocks will be improved, probably by a greater factor than the astronomical determinations," which in retrospect was a considerable understatement [30]. To Bullard and many others, it was already clear when the ephemeris second became the SI second that atomic clocks represented the future of timekeeping.

In June 1955, due in part to the fact that there were no reliable atomic clocks then operating in the United States, NPL and USNO began collaborating on a program in which the goal was to determine the frequency of cesium with respect to the ephemeris second [31, 32]. It was a worthy goal, because redefining the second based on an atomic transition would allow large numbers of laboratories to simultaneously produce a physical realization of the second, which could in turn be used as the reference for other laboratory measurements. To achieve the goal, NPL and the USNO would need to compare their clocks and measure their differences, which they did from mid-1955 until the end of the first quarter of 1958, a period of 2.75 years. The USNO clock was based on a 2.5 MHz quartz oscillator manufactured by Western Electric,¹ which was steered to ephemeris time by applying corrections obtained with the dual-rate Moon camera. The NPL clock was based on their new cesium standard, which was now considered to be accurate to within $5 \cdot 10^{-10}$. Because the two clocks were on opposite sides of the Atlantic Ocean, the comparison could only be made by simultaneously comparing both clocks to radio signals that could be received at both laboratories, a measurement technique now known as common-view time transfer. Several time signal broadcast stations were involved in the measurement, including NBS radio station WWV,

then located in Beltsville, Maryland, in the United States, and radio stations MSF and GBR, located in the United Kingdom [31].

To make the uncertainty of the measurement results as small as possible, four different solutions were calculated to determine the effects of using different sets of data. The final measurement result was the average of the four solutions, and it was published as 9 192 631 770 cycles/s in August 1958, with an estimated measurement uncertainty of ± 20 cycles/s. The measurement result applied to cesium in a zero-magnetic field and to ephemeris time at the beginning of 1957 [32]. The measurement uncertainty was limited not by the cesium standard, but rather due to the many complexities involved in measuring ephemeris time. Multiple sources of short-term noise, including the radio signals used in common-view time transfer and large uncertainties in the Moon camera itself, meant that the necessary resolution could only be achieved with very long averaging periods of at least 100 days. If we stop to consider the limitations of the technology that existed in the 1950s and the large number of complex steps that were required, the measurement results were remarkably accurate. As Leschiutta wrote in 2005:

"As a personal remark, taking into account the capabilities of the timing emissions at the moment, of the frequency standards available, of the inevitable scatter of the moon camera, and some other factors, not least the widespread use and abuse in 'touching' the piezo-oscillator, it is almost impossible to explain the accuracy of the Markowitz determination. Similar events, i.e. results surpassing the capabilities of the moment, are not uncommon in the history of science that sometimes is prone to accepting the intervention of a serendipity principle. The other possible explanation calls for a first class understanding of physics, coupled with scientific integrity." [31]

The publication of the Markowitz/Essex measurement results [32] made it inevitable that an atomic time scale would eventually supplant the existing astronomical time scales [33, 34], yet still nearly a decade passed before the definition of the second was changed. At least part of the delay was because atomic standards other than cesium devices, specifically those based on hydrogen and thallium, were being considered as the basis for the redefinition. Jean Terrien, the director of the Bureau International des Poids et Mesures (BIPM), wrote in August 1967 that the definition had nearly been changed in 1964, but

"it was considered premature in 1964 to redefine the second in terms of the caesium frequency because the hydrogen maser was expected to offer, within a short while, a better base for such a redefinition."

¹ Certain commercial equipment, instruments, or materials are identified in this paper for the purposes of illustration and to provide a historical narrative. Such identification does not imply recommendation or endorsement by the National Institute of Standards and Technology, nor does it imply that the materials or equipment identified are necessarily the best available for the purpose.

Contrary to expectation, it happened that the hydrogen maser, though potentially slightly superior to the caesium beam frequency standard, progressed only slowly (one of its uncertainties comes from the wall shift), and that caesium standards were so much improved that an almost perfect confidence was acquired in their reliability, within the most severe present needs of precision. Though it was agreed that the improvement of other atomic frequency standards, using for example hydrogen or thallium atoms, ought to be pursued actively, the Comité Consultatif pour la Définition de la Seconde, at its meeting of 12–13 July 1967, recommended unanimously to the Comité International des Poids et Mesures not to wait any longer ...” [35]

Thus, in October 1967, the SI second was formally redefined by the Comité International des Poids et Mesures (CIPM) as:

“the duration of 9 192 631 770 periods of the radiation corresponding to the transition between the two hyperfine levels of the ground state of the caesium 133 atom.” [1]

As of 2017, the definition of the SI second remains the same, except for a slight amendment made in 1997. Calculations made by Wayne Itano of NBS in the early 1980s [36] revealed that black-body radiation can cause noticeable frequency shifts in cesium standards, and his work eventually resulted in an addendum to the definition. The CIPM affirmed in 1997 that the definition refers

to “a cesium atom at rest at a thermodynamic temperature of 0 K.” This meant that an optimal realization of the SI second requires the cesium atoms to be in a field-free environment where the temperature is absolute zero and where the atoms have no residual velocity.

When the atomic second became the SI second 1967, it was of course known that atomic time would continuously diverge from astronomical time, so periodic corrections would be needed to keep the new atomic time scale, called Coordinated Universal Time (UTC), in step with astronomical time. The term UTC had been used informally since the early 1960s, as the United Kingdom and United States had begun coordinating frequency adjustments made to their time scales beginning on January 1, 1960. Other countries followed suit, and in 1961, the Bureau International de l’Heure in Paris, France, began to oversee the international time coordination process [28]. Frequency adjustments continued to be used for almost five years after the new SI definition, but the procedure changed when the first leap second was inserted on June 30, 1972. Leap seconds are periodically added to UTC to keep it within ± 0.9 s of UT1, an astronomical time scale based on the mean solar second. Adding a leap second to UTC stops atomic time for one second to allow astronomical time to catch up.

The exact divergence between astronomical and atomic time is difficult to model or predict, but there are two general reasons why leap seconds are periodically needed. Reason one can be traced to the original Markowitz/Essen measurement; the atomic second was defined with respect to the ephemeris second, which then served as the SI second, and not with respect to the mean solar second, which forms the basis for UT1. Ephemeris seconds were thus already slightly shorter than mean solar seconds, and this characteristic was passed along to the atomic second, making the need for periodic leap seconds inevitable. However, even if the atomic second had been defined with respect to the mean solar second, there would still be a need for leap seconds, albeit a much smaller number. Leap seconds would still have been occasionally needed due to reason two, which is simply that the mean solar second is gradually getting longer because Earth’s average rotational rate is gradually slowing down. The second reason is commonly cited as the chief reason for leap seconds, but it has had far less impact on the number of leap seconds that have occurred so far than the first reason.

6. Commercial Cesium Clocks

Commercial cesium clocks were available even before the publication of the Markowitz/Essen

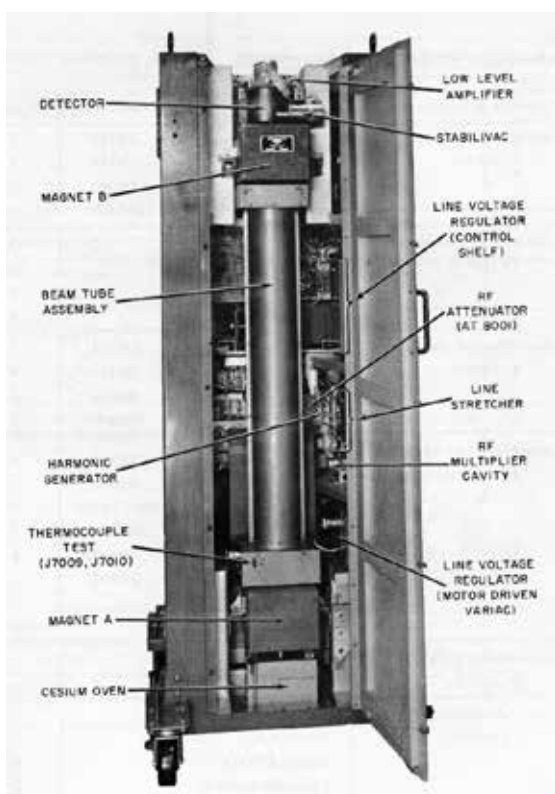


Figure 8: Interior of Atomichron NC2001, showing vertical cesium beam tube (courtesy of Tom Van Baak, leapsecond.com).

measurement in 1958, and several different models designed and manufactured in the United States were in use in laboratories prior to the adoption of the cesium definition in 1967. The availability of these commercial devices, as well as the measurement results and experience gained from their use, undoubtedly expedited the acceptance of the atomic definition of the SI second.

The first commercial cesium clock was introduced in October 1956, just slightly more than one year after the introduction of the NPL clock. Named the Atomichron, the clock was manufactured by the National Company of Malden, Massachusetts, and developed by a team led by Jerrold Zacharias of MIT [8, 37], an experienced researcher who had previously collaborated on early molecular beam experiments with Rabi [4] at Columbia University. About 50 Atomichrons were sold between 1956 and 1960, most of them to the U.S. military, with the first unit being delivered to the Naval Research Laboratory in Washington, D.C., in September 1956. The original models, numbers NC1001 and NC2001, were large devices, housed in an equipment rack that was about 0.6 m deep by 0.6 m wide, and about 2.1 m tall, and they weighed more than 350 kg. The cesium beam tube was vertically mounted (Fig. 8) and nearly as tall as the equipment rack [8, 38]. The resonance frequency of the first model was 9 192 631 830 cycles/s (± 10 cycle/s), or 60 Hz higher than the yet-to-be published definition [37].

At least two Atomichrons were in operation at NBS in 1958–1959, where they were compared to the NBS standard to within parts in 10^{10} [23]. An additional two devices were transported to the United Kingdom and compared to Essen's clock at NPL with similar results [39]. Much smaller Atomichrons appeared later, produced mostly for the U.S. Air Force. However, the National Company ultimately failed due to financial difficulties and sold its assets, intellectual property, and trade names to Frequency Electronics, Inc., in 1969 [8, 38].

A far more successful line of cesium clocks was developed through the efforts of Leonard Cutler and his colleagues at the Hewlett-Packard Company in Santa Clara, California. Born in Los Angeles, California, in 1928, Cutler earned a Ph.D. from Stanford University and spent over four decades working for Hewlett-Packard and its successor, Agilent Technologies, where he designed or co-designed numerous frequency standards and clocks. During Cutler's tenure, Hewlett-Packard produced a series of cesium clocks that were reliable enough to run continuously for years and small enough to fit into standard equipment racks and were thus introduced into many calibration and metrology laboratories in the 1960s. The first of these clocks was the Hewlett-Packard 5060,

introduced in 1964, some three years before the redefinition of the SI second (Fig. 9). The first cesium standard to feature all solid-state electronics, the length of its Ramsey cavity was just 12.4 cm, and the device weighed less than 30 kg [38, 40]. By 1966, the 5060 had a specified accuracy of $1 \cdot 10^{-11}$ [21]. The 5060 was followed by the 5061, manufactured from 1967 until the early 1990s, and then by the 5071, which first appeared in the Hewlett-Packard catalog in 1993. With an internal microprocessor and an improved cesium beam tube, the 5071 was more stable and accurate than all its predecessors [38]. By 1999, it had a specified accuracy of about 50 ns/day ($5 \cdot 10^{-13}$), and could easily be adjusted to keep time within a few nanoseconds per day (parts in 10^{14}). Later manufactured by Agilent and Symmetricom and now manufactured by Microsemi, the 5071 continues to serve as the primary frequency and time standard at numerous laboratories.

7. Principles of a Cesium Clock

The designs of cesium beam standards can vary significantly, but all designs, including those of the commercially available clocks still being sold in 2017 and the early cesium beam clocks built at NBS, NPL, and other laboratories, are based on principles that can be traced back to the seminal work of Rabi and Ramsey.

To understand how a cesium clock implements the SI definition of the second, consider that cesium is a complicated atom with $F = 3$ and $F = 4$ ground states (Fig. 10). Each atomic state is characterized not only by the quantum number F , but also by a second quantum number, m_F , which can have integer values between $-F$ and $+F$.

The splitting of the $F = 3$ and $F = 4$ states into the various m_F sublevels occurs in the presence of a magnetic field. There are 16 possible m_F sublevels

Figure 9:
Hewlett-Packard
5060A cesium beam
atomic clock with
cover removed
showing beam tube
on left (courtesy
of Tom Van Baak,
leapsecond.com).



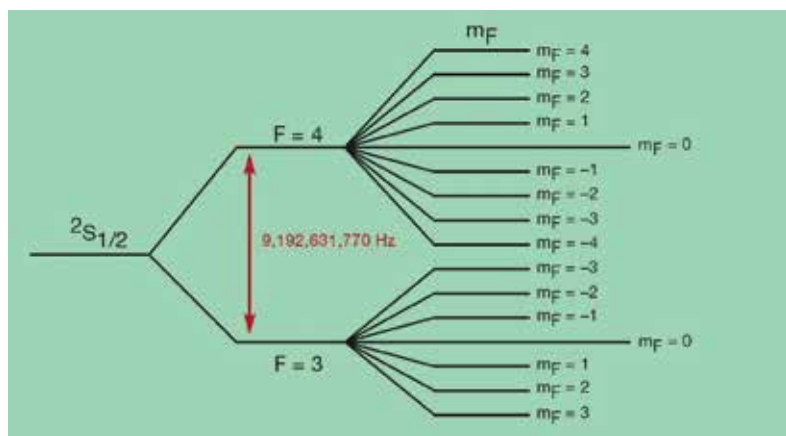


Figure 10:
The cesium clock
transition.

in the ground state of cesium, but the frequency of the $|4,0\rangle \leftrightarrow |3,0\rangle$ hyperfine transition was the one chosen to define the second. The $|4,0\rangle \leftrightarrow |3,0\rangle$ transition had several advantages that made it the best choice for clocks. There was a very low probability of a spontaneous transition occurring during the observation time. The transition was also relatively easy to detect with electronic systems that were already available when Rabi and others began their experiments. Perhaps most importantly, the transition was relatively insensitive to magnetic fields, so small magnetic fields near a cesium clock would have little effect on its frequency.

Figure 11 provides a simplified schematic of a cesium beam clock. As shown on the left side of Fig. 11, ^{133}Cs atoms are heated to a gaseous state in an oven. A beam of atoms emerges from the oven at a temperature near 100 °C and travels through a spatially varying magnetic field, where the beam is split into two atomic beams with different magnetic states that follow different trajectories. One beam is absorbed by the getter and is of no further interest. The other beam is deflected into the microwave interrogation cavity (i.e., Ramsey cavity), which exposes the atoms to microwaves in two spatially separate pulses.

While inside the Ramsey cavity, the cesium beam is exposed to a microwave signal. This signal is generated by a frequency synthesizer driven

by a quartz oscillator. If this frequency exactly equals the cesium resonance frequency, the atoms will change their magnetic state (the desired $|4,0\rangle \leftrightarrow |3,0\rangle$ energy transition shown in Fig. 10) while inside of the Ramsey cavity. After leaving the cavity, the atoms pass through a second spatially varying magnetic field. These magnets direct only the atoms that changed state to the detector; the other atoms are directed to a getter and absorbed. The detector sends a feedback signal to a servo circuit that continually tunes the quartz oscillator so that the maximum number of atoms reaches the detector, thereby locking the frequency of the microwaves to the cesium hyperfine transition frequency. The process of keeping the quartz oscillator locked to the cesium frequency is analogous to carefully and continuously tuning a radio dial to ensure that the loudest and clearest signal is always heard [41, 42].

8. Summary and Conclusion

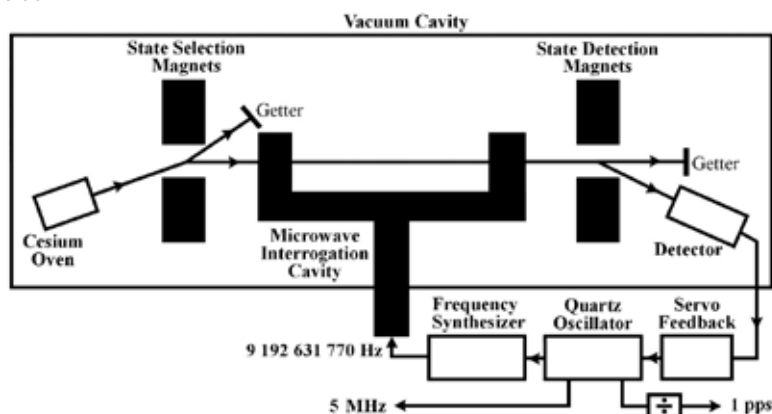
Countless contributions leading to the atomic definition of the SI second in 1967 were made by U.S. scientists and engineers, and so any historical account of this length cannot be comprehensive. The names of many individuals and organizations who helped to establish the modern era of atomic timekeeping have been regrettably omitted. This historical overview has briefly described the technology and noted the accomplishments of individuals, including Isidor Isaac Rabi, Sidney Millman, Polykarp Kusch, Norman Ramsey, Jerrold Zacharias, Harold Lyons, Jesse Sherwood, Richard Mockler, William Markowitz, and Leonard Cutler, and of organizations, including Columbia University, the Massachusetts Institute of Technology, Harvard University, the National Bureau of Standards, the U.S. Naval Observatory, the National Company, and the Hewlett-Packard Company. Their contributions, and the societal benefits of atomic timekeeping that we now enjoy because of their work, should not be forgotten.

Acknowledgments

The author thanks the internal reviewers at NIST, Elizabeth Donley and Tom Heavner, and the external peer reviewers, for many helpful comments and corrections.

This paper is a contribution of the United States Government, and is not subject to copyright. Commercial products and companies are identified for the purposes of illustration and description and to provide a historical narrative. This does not imply endorsement by NIST.

Figure 11:
Block diagram of
a cesium beam
clock.



References

- [1] Resolution 1 of the 13th Conference Generale des Poids et Mesures (CGPM), 1967
- [2] W. Thomson and P. Tait, "*Elements of Natural Philosophy*," Cambridge: At the University Press, pp. 61–62, 1879
- [3] W. Snyder, "Lord Kelvin on Atoms as Fundamental Natural Standards (for Base Units)," IEEE Transactions on Instrumentation and Measurement, vol. **IM-22**, p. 99, March 1973.
- [4] I. Rabi, J. Zacharias, S. Millman, and P. Kusch, "A New Method of Measuring Nuclear Magnetic Moment," *Physical Review*, vol. **53**, p. 318, February 1938
- [5] N. Ramsey, "*Molecular Beams*," Clarendon Press, Oxford, 1956
- [6] N. Ramsey, "History of Atomic Clocks," *Journal of Research of the National Bureau of Standards*, vol. **88**, pp. 301–320, September-October 1983
- [7] W. Laurence, "Cosmic Pendulum for Clock Planned," *New York Times*, p. 34, January 21, 1945
- [8] P. Forman, "Atomichron™: The Atomic Clock from Concept to Commercial Product," *Proceedings of the IEEE*, vol. **73**, pp. 1181–1204, July 1985
- [9] S. Millman and P. Kusch, "On the Radiofrequency Spectra of Sodium, Rubidium, and Caesium," *Physical Review*, vol. **58**, pp. 438–445, September 1940
- [10] P. Forman, "The first atomic clock program: NBS, 1947–1954," *Proceedings of the 1985 Precise Time and Time Interval Meeting (PTTI)*, pp. 1–17, 1985
- [11] National Bureau of Standards, "The Atomic Clock: An Atomic Standard of Frequency and Time," NBS Technical Report 1320, p. 27, January 1949
- [12] H. Lyons, "The Atomic Clock," *Instruments*, vol. **22**, pp. 133–135, December 1949
- [13] H. Lyons, "Spectral lines as frequency standards," *Annals of the New York Academy of Sciences*, vol. **55**, pp. 831–871, November 1952
- [14] H. Lyons and B. Huston, "Atomic Clock," United States Patent 2,699,503 (Applied for April 30, 1949, granted January 11, 1955)
- [15] H. Lyons, "Microwave Frequency Dividers," *Journal of Applied Physics*, vol. **21**, pp. 59–60, January 1950.
- [16] H. Lyons, "The Atomic Clock: A Universal Standard of Frequency and Time," *American Scholar*, vol. **19**, pp. 159–168, April 1950
- [17] J. Sherwood, H. Lyons, R. McCracken, and P. Kusch, "High frequency lines in the hfs spectrum of cesium," *Bulletin of the American Physical Society*, vol. **27**, p. 43, 1952
- [18] N. Ramsey, "A Molecular Beam Resonance Method with Separated Oscillating Fields," *Physical Review*, vol. **78**, pp. 695–699, June 1950
- [19] H. Lyons, "Spectral lines as frequency standards," National Bureau of Standards Report 1848, p. 71, August 8, 1952
- [20] W. Snyder and C. Bragaw, "*Achievement in Radio: Seventy Years of Radio Science, Technology, Standards, and Measurement at the National Bureau of Standards*," National Bureau of Standards Special Publication 555, October 1986
- [21] R. Beehler, "A Historical Review of Atomic Frequency Standards," *Proceedings of the IEEE*, vol. **55**, pp. 792–805, June 1967
- [22] R. Mockler, R. Beehler, and C. Snider, "Atomic Beam Frequency Standards," *IRE Transactions on Instrumentation*, vol. **I-9**, pp. 120–132, September 1960
- [23] R. Mockler, R. Beehler, and J. Barnes, "An evaluation of a cesium beam frequency standard," National Bureau of Standards Report 6075, October 1959
- [24] R. Beehler and D. Glaze, "The Performance and Capability of Cesium Beam Frequency Standards at the National Bureau of Standards," *IEEE Transactions on Instrumentation and Measurement*, vol. **15**, pp. 48–55, March-June 1966
- [25] D. Henderson, "Essen and the National Physical Laboratory's atomic clock," *Metrologia*, vol. **42**, pp. S4–S9, June 2005
- [26] L. Essen and J. Parry, "An Atomic Standard of Frequency and Time Interval: A Caesium Resonator," *Nature*, vol. **176**, pp. 280–281, August 1955
- [27] C. Audoin and B. Guinot, "*The Measurement of Time: Time, Frequency and the Atomic Clock*," Cambridge University Press, 2001
- [28] D. McCarthy and P. Seidelmann, "*Time: From Earth Rotation to Atomic Physics*," Wiley-VCH, 2009
- [29] W. Markowitz, "Photographic Determination of the Moon's Position, and Applications to the Measure of Time, Rotation of the Earth, and Geodesy," *The Astronomical Journal*, vol. **59**, pp. 69–73, March 1954
- [30] E. Bullard, "Definition of the Second of Time," *Nature*, vol. **176**, p. 282, August 1955
- [31] S. Leschiutta, "The definition of the 'atomic' second," *Metrologia*, vol. **42**, pp. S10–S19, June 2005
- [32] W. Markowitz, R. Glenn Hall, L. Essen and J. Parry, "Frequency of Cesium in Terms of Ephemeris Time," *Physical Review Letters*, vol. **1**, no. 3, pp. 105–107, August 1958
- [33] W. Markowitz, "The Atomic Time Scale," *IRE Transactions on Instrumentation*, vol. **I-11**, no. 3, pp. 239–242, December 1962
- [34] R. Tipson, "Sun-time Replaced by Atomic Clocks," *The Capital Chemist*, vol. **11**, pp. 255–256, November 1961
- [35] J. Terrien, "A major step towards the redefinition of the second, the unit of time in the International System of Units," *Metrologia*, vol. **3**, no. 4, p. 130, October 1967

- [36] W. Itano, L. Lewis, D. Wineland, “Shift of $^2S_{1/2}$ hyperfine splittings due to blackbody radiation”, *Physical Review A*, vol. **25**, pp. 1233–1235, February 1982
- [37] “Atomichron – World’s Most Accurate Clock”, *Radio and Television News*, p. 63, January 1957
- [38] L. Cutler, “Fifty years of commercial caesium clocks”, *Metrologia*, vol. **42**, S90–S99, June 2005
- [39] A. McCoubrey, “Results of comparison: Atomichron-British cesium beam standard”, *Proceedings of the 12th Annual Frequency Control Symposium*, pp. 648–664, 1958
- [40] A. Bagley and L. Cutler, “A Modern Solid-State Portable Cesium Beam Frequency Standard”, *Proceedings of the 18th Annual Frequency Control Symposium*, pp. 344–365, 1964
- [41] W. Itano and N. Ramsey, “Accurate Measurement of Time”, *Scientific American*, vol. **269**, pp. 56–65, July 1993
- [42] F. Major, “*The Quantum Beat: The Physical Principles of Atomic Clocks*”, Springer-Verlag, New York, 1998

Atomuhren und Navigation mit Satelliten-Systemen (GNSS)

Gerhard Beutler*

Präambel

Im Jahr 2017 wird der fünfzigste Jahrestag der derzeit gültigen Definition der SI-Sekunde gefeiert. Das Ereignis ist der Anlass zum vorliegenden Band der Mitteilungen der Physikalisch-Technischen Bundesanstalt Braunschweig und Berlin – und damit auch zu diesem Beitrag.

Das Ereignis sollte im Kontext von drei anderen Entwicklungen gesehen werden: Dem ersten künstlichen Erdsatelliten „Sputnik 1“, dem ersten Satelliten-Navigationssystem „Transit“ und der Very Long Baseline Interferometry (VLBI).

Der Start des Erdsatelliten Sputnik 1 (Bild 1) am 4. Oktober 1957 markiert den Beginn der Raumfahrtära. Somit feiert man im Jahr 2017 auch sechzig Jahre Raumfahrt. Im Jahr 1967 wurde das Satelliten-Navigationssystem Transit für zivile Anwendungen geöffnet. Da sein Nachfolger GPS

(Global Positioning System) von Anfang an ein dualuse System war, feiert man 2017 also auch 50 Jahre Satelliten-Positionierung im Dienste von Wissenschaft und Gesellschaft. Im Jahr 2017 wird weiter der fünfzigste Geburtstag von Very Long Baseline Interferometry (VLBI) gefeiert: Erste brauchbare Messungen mit VLBI-Teleskopen wurden im Frühjahr und Frühsommer 1967 in Kanada und unabhängig davon in den USA durchgeführt. VLBI, die interferometrische Technik par excellence, enthält wichtige Elemente der heutigen Globalen Navigations-Satellitensysteme (GNSS). Die vier Jahrestage sind zwar unabhängig voneinander und sie werden wohl auch unabhängig voneinander gewürdigt, es gibt aber Berührungspunkte. Tabelle 1 enthält, neben den bereits erwähnten, weitere in unserem Zusammenhang wichtige Ereignisse der letzten sechzig Jahre.

* Prof. Dr. Gerhard Beutler, Astronomisches Institut der Universität Bern, gerhard.beutler@aiub.unibe.ch

1957	Start von Sputnik 1, Beginn der Raumfahrtära
1959	Start des ersten Transit-Satelliten
1964	Transit als operationell erklärt
1967	Transit für zivile Anwendungen geöffnet
1967	Erste Definition der SI-Sekunde
1967	Erste erfolgreiche VLBI-Beobachtungen
1978	Start des ersten GPS-Satelliten
1982	Start der ersten GLONASS-Satelliten
1990	Einschalten von Selective Availability (S/A) bei GPS
1992	International-GPS-Service-Test-Kampagne, gefolgt von Pilot-Dienst
1993	GLONASS operationell erklärt
1994	Offizieller Beginn des IGS als wissenschaftlicher Dienst der Internationalen Assoziation für Geodäsie
1995	GPS als operationell erklärt
1995	Ende Dezember: Offizielles Ende von Transit als Positionierungssystem
2000	2. Mai 2000: Ausschalten von S/A bei GPS
2005	Start von GIOVE-A (erster Galileo-Test-Satellit)
2011	GLONASS-Konstellation erneut vollständig verfügbar (erneut operationell)
2011	Beidou-2 mit 5 GEO-, 5 IGSO- und 4 MEO-Satelliten operationell

Tabelle 1: Ereignisse im Bereich Globale Navigation mit Satellitensystemen (GNSS)

Das Transit System wurde auch „Navy Navigation Satellite System (NNSS)“ oder „NAVSAT“ genannt, GPS wurde offiziell als „NAVSTAR GPS (Navigational Satellite Timing and Ranging Global Positioning System)“ bezeichnet. Bei „Selective Availability“ handelt es sich um eine künstliche Verschlechterung der Satellitenuhren.

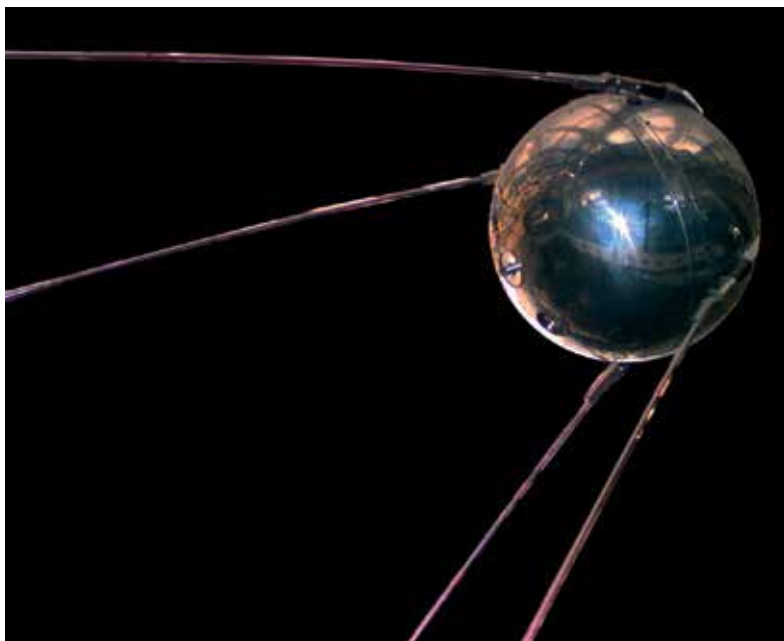
Nicht in die Tabelle aufgenommen wurden regionale Systeme wie QZSS (japanisches Quasi-Zenit-Satellitensystem) und IRNSS (Indisches Regionales Satellitensystem).

Sputnik 1 und das Transit-Doppler-System

Der erste künstliche Erdsatellit Sputnik 1, dargestellt in Bild 1, verfügte über einen Kurzwellensender von 1 Watt Leistung. Das „Herz“ des Satelliten bestand aus einem quartzesteuerten Oszillator. Forscher und Funkamateure konnten die bedeutenden Dopplerverschiebungen der von Sputnik 1 ausgesandten Signale während mehrerer Wochen messen, aufzeichnen und analysieren.

Glaubt man den Referenzen [1, 2], einer sehr schönen und informativen Geschichte des US-amerikanischen GPS, hat die Messung des Dopplereffektes von Sputnik 1 und die Bahnbestimmung mithilfe dieser Messungen direkt zur Entwicklung des ersten globalen Positionierungssystems Transit geführt. Wichtig war die Idee, die Problemstellung umzukehren: Bei bekannter Satellitenbahn sollte es möglich sein, aus einer kurzen Sequenz von Dopplermessungen, gemessen mit einem Radio-ähnlichen Empfänger auf der Erdoberfläche, die Position dieses Empfängers zu bestimmen. Das von 1964 bis 1995 operationelle US-amerikanische Transit-Satellitensystem wies schon viele der noch heute wesentlichen Elemente der globalen Navigation mit Satelliten auf:

Bild 1:
Sputnik 1 (NASA -
<http://nssdc.gsfc.nasa.gov/database/MasterCatalog>)



1. Zur globalen Navigation benötigt man nicht *einen* Satelliten, sondern *ein System* baugleicher Satelliten, deren wichtigste Komponente ein hochpräziser Oszillator ist. Von diesem Oszillator werden die Trägerwellen der abgestrahlten Signale abgeleitet.
2. Um eine quasi-gleichzeitige Positionierung von allen Punkten der Erdoberfläche aus zu ermöglichen, dürfen die Satelliten nicht in einer Bahnebene sein, sie müssen vielmehr auf mehrere Bahnebenen verteilt werden.
3. Auf See kann das Problem bei Transit auf ein zweidimensionales reduziert werden, da die Höhe (Meereshöhe) bekannt ist.
4. Allwettertauglichkeit bedingt Satellitensignale im Radio- oder Mikrowellenbereich.
5. Zusätzlich zur eigentlichen Messgröße, der Dopplerverschiebung der Trägerwellen, muss der Satellit dem Nutzer seine Bahn sowie Stand und Gang der Satellitenuhr relativ zu einer Systemzeit in Echtzeit bekanntgeben.
6. Zur Elimination der ionosphärischen Refraktion im Radio- oder Mikrowellenbereich müssen die Satellitensignale auf mindestens zwei gut separierten Trägerfrequenzen übertragen werden.

War das Transit-System nach heutigem Verständnis wirklich ein *globales* Satelliten-Navigationssystem?

1. Mit fünf Satelliten in fünf um je 72° im Äquator gegeneinander verdrehten Bahnebenen einer Neigung von $i \approx 90^\circ$ gegenüber der Äquatorebene war Transit effektiv ein globales System.
2. Wartete man an einem beliebig gewählten Punkt der Erdoberfläche lange genug (maximal bis zu etwa hundert Minuten am Äquator), konnte man mit Sicherheit den Vorbeiflug eines Transit-Satelliten beobachten: mit einer Bahnhöhe von etwa 1075 km war es am Äquator nicht möglich, von einem Punkt aus mehr als einen Satelliten gleichzeitig zu beobachten. Quasi-gleichzeitige Positionierung bedeutete daher bei Transit „innerhalb von etwa zwei Stunden (die Umlaufzeit der Satelliten betrug etwa 106 Minuten)“.
3. Unter Nutzung nur eines Satellitendurchganges war die „genaue“ Positionierung in Echtzeit zweidimensional, wobei die Genauigkeit je nach Geometrie und je nach Qualität der Satelliten-Ephemeride bei etwa 10–100 m, im Mittel bei etwa 25 m lag.
4. Transit sendete Signale auf Trägerwellen mit Wellenlängen im Dezimeter- bis Meterbereich, also im hochfrequenten UKW-Bereich aus. Die Wetterunabhängigkeit war damit gegeben.

5. Die Bahnen der Transit-Satelliten wurden aus den Messungen von permanenten Bodenstationen, dem so genannten *Ground Segment* von Transit, berechnet. Die Bahnprognosen wurden den Satelliten von der Erde aus regelmäßig, im Abstand von etwa 12 Stunden, mitgeteilt. Die Satelliten haben diese Prognosen als *broadcast ephemeris* dem Signal aufmoduliert und dem Nutzer zur Verfügung gestellt. Das von Transit entwickelte Prozedere der Bahnbestimmung und -prognose wurde im Wesentlichen vom Nachfolge-System GPS übernommen.
6. Transit sendete die Signale auf Frequenzen von 400 MHz und 150 MHz, was Wellenlängen von etwa 0,75 m respektive 2 m entspricht. Dadurch war es möglich, die ionosphärischen Laufzeitunterschiede der Signale (ionosphärische Refraktion) in guter Näherung zu eliminieren. Wegen der (fast) polaren Bahnen eignete sich Transit auch sehr gut zum Studium der Ionosphäre: Innerhalb von „nur“ zwei Stunden wurde die ganze Erde in zehn um 36° gegeneinander versetzten Meridianen abgetastet. Daher konnte das System, nach seinem offiziellen Ende als Navigationssystem im Dezember 1995, weiter zur Untersuchung der Ionosphäre der Erde genutzt werden.

Bild 2 zeigt einen Transit-Satelliten. Die Ausrichtung der Sendeantenne zur Erde wurde durch Gravitations-Stabilisierung, realisiert durch einen langen, von der Erde aus gesehen nach oben ausgerichteten Mast, mit einem Gegengewicht an dessen Ende, erreicht. Die Transit-Satelliten hatten in den frühen Jahren eine Masse von 50–60 kg, spätere Satelliten hatten eine etwas größere Masse. Die zum Betrieb der Satelliten notwendige Energie wurde bei den operationellen Transit-Satelliten mit Sonnenpaneelen erzeugt.

Das im Jahr 1967 vom damaligen Vizepräsidenten der USA ausgesprochene Angebot, Transit auch zivil zu nutzen, wurde von der Wissenschaft, nach anfänglichen Berührungsängsten, mit Begeisterung aufgenommen. Unzählige Doppler-Kampagnen, zum Teil auch in vorher schlecht vermessenen Teilen der Erde, wie z. B. in Afrika, sind noch heute von Bedeutung. Das *World Geodetic System 72* (WGS72) beruht zu einem großen Teil auf den Dopplermessungen erdfester Stationen zu den Transit-Satelliten. Das WGS72 ist der Vorläufer des WGS84, welches wiederum der Vorläufer der heute allgemein akzeptierten globalen Referenzsysteme ist, die vom *International Earth Rotation and Reference Systems Service* (IERS) unter Verwendung aller modernen Beobachtungsmethoden herausgegeben werden. Sechs internationale wissenschaft-

liche Symposien der Reihe *International Geodetic Symposia on Satellite Positioning* geben Zeugnis von der Bedeutung von Transit für Wissenschaft und Gesellschaft. Im Vorwort der *Proceedings* des sechsten und letzten dieser Symposien in Ohio im Jahr 1992 [3] liest man allerdings „*the symposium was the sixth in a series which began in 1976 in Las Cruces, New Mexico. The period from the first symposium to the present represented a complete transition from an all Navy Navigation Satellite System (NNSS) to an all Global Positioning System (GPS) technical program.*“ Es ist aus heutiger Sicht schwer zu verstehen, dass Transit, noch bis Ende 1995 voll operationell, am Symposium von 1992 nicht einmal durch einen Statusbericht vertreten war. Offenbar wurde schon damals die Tagesaktualität höher als historische Korrektheit gewichtet.

Während die Echtzeit-Positionierung von Transit wie erwähnt im Genauigkeitsbereich von 10 bis 100 Metern lagen, erzielte man in wissenschaftlichen Anwendungen durch die Kombination vieler Satellitendurchgänge durchaus schon Genauigkeiten im Meter-Bereich. Das Programmsystem GEODOP, entwickelt von Jan Kouba, wurde zur Auswertung sehr vieler nationaler und internationaler Doppler-Kampagnen verwendet [4]. Viele davon wurden organisiert vom deutschen Institut für Angewandte Geodäsie (IfAG) in Frankfurt, dem Vorgänger des heutigen Bundesamtes für Kartographie und Geodäsie (BKG). Transit-Empfänger konnten auch zur Zeitsynchronisation lokaler Uhren genutzt werden, wobei eine Genauigkeit von etwa 10 μ s erreicht wurde.

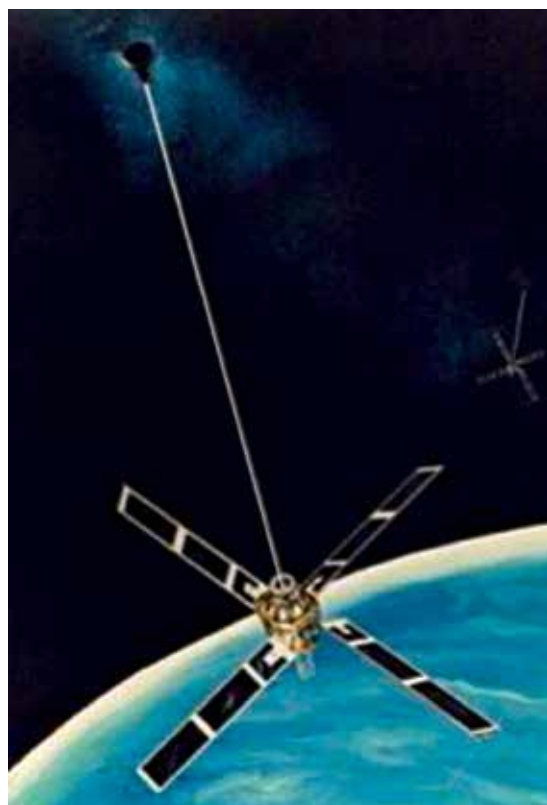
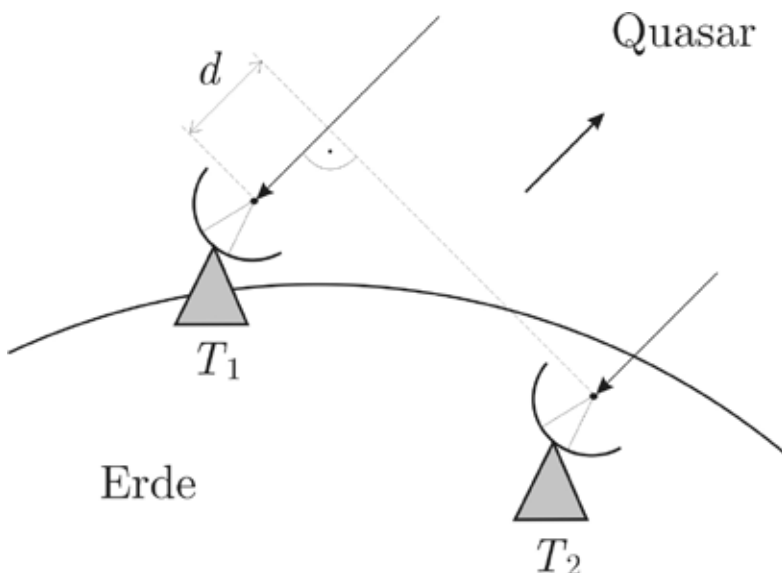


Bild 2:
Transit-Satellit (<https://commons.wikimedia.org/w/index.php?curid=172629>)

VLBI – die Jagd nach dem Millimeter und der Mikrobogensekunde (μs)

VLBI-Teleskope zeichnen die Mikrowellen-Signale von Quasaren auf. Quasare sind sehr weit entfernte Radiogalaxien. Das Prinzip wird durch Bild 3 veranschaulicht: Die Mikrowellen-Signale eines Quasars werden von den Radioteleskopen T_1 und T_2 gemeinsam mit der Atomzeit der Stationen (realisiert durch Wasserstoffmaser), aufgezeichnet. Eigentliche Observable ist die vom Quasar aus gemessene Distanzdifferenz d zwischen den Teleskopen T_1 und T_2 , die aus den Differenzen der Ankunftszeiten ein und derselben (ebenen) Wellenfronten bei T_1 und T_2 , nach Multiplikation mit der Lichtgeschwindigkeit c mit Millimeter-Genauigkeit rekonstruiert werden kann. Voraussetzung ist eine sehr genaue Synchronisation der Wasserstoffmaser im VLBI-Netz. Die Methode stellte in den 1960er-Jahren eine Revolution dar. Mit VLBI gelang es, ein zwar nicht sehr dichtes Netz von Fundamentalstationen auf der Erde mit Millimeter-Genauigkeit festzulegen. Da die Teleskope mit der Erde rotieren, die Quasare aber fest mit dem Inertialsystem verknüpft sind, ermöglicht es VLBI, die Koordinaten des Rotationspols der Erde auf der Erdoberfläche, die Differenz UT1–UTC, also den Unterschied zwischen den Zeitskalen, definiert durch Erdrotation bzw. Atomzeit, sowie die Nutationswinkel, welche die Lage der Rotationsachse in einem inertialen Bezugssystem definieren, zu bestimmen – mit einer damals und auch heute als traumhaft empfundenen Genauigkeit. Mit einem Schlag kam man vom Meter- in den Millimeter-Bereich und vom Millibogensekunden-(ms)- in den Mikrobogensekunden-(μs)-Bereich. Das sind geradezu unheimlich anmutende Genauigkeiten, entspricht doch eine Millibogensekunde auf der Erdoberfläche 31 mm, eine Mikrobogensekunde 0,03 mm!

Bild 3:
Vermessung mit
VLBI (aus [5])



In Bild 3 könnte (1) der Quasar durch einen Satelliten ersetzt werden (wobei die Wellenfronten dann nicht mehr eben, sondern kugelförmig wären), (2) die Radioteleskope durch Satelliten-Empfänger respektive durch deren Empfangsantennen und (3) die Atomuhren durch die Kristalloszillatoren der Empfänger. Die erreichbare Genauigkeit wäre wegen der bescheidenen Stabilität der Empfänger-Oszillatoren nicht sehr hoch. Natürlich könnte dieser Nachteil durch bessere Empfängeruhren ausgeglichen werden, allerdings zu einem sehr hohen Preis (Vorschläge in der Richtung hat es gegeben). Moderne GNSS bieten eine weit intelligentere und günstigere Lösung dieses Problems, der wir uns jetzt zuwenden.

Globale Navigation mit modernen GNSS

Das amerikanische GPS und das russische GLONASS waren die ersten globalen Navigations-Satellitensysteme. Alle globalen Navigations-Satellitensysteme werden mit GNSS abgekürzt. Die ersten Satelliten wurden in den Jahren 1978 (GPS) und 1982 (GLONASS) in Umlaufbahnen gebracht, die Systeme wurden 1993 (GLONASS) und 1995 (GPS) von den beiden Betreibern als voll funktionsfähig oder operationell erklärt. Während GPS in zivilen Anwendungen bereits ab Beginn der 1980er-Jahre stark beachtet und genutzt wurde, wurde GLONASS außerhalb der damaligen Sowjetunion kaum beachtet. Das änderte sich erst in der zweiten Hälfte der 1990er-Jahre, als (1) kommerzielle GPS-/GLONASS-Empfänger produziert wurden und als (2) der IGS (siehe nächste Seite) es sich zum Ziel setzte, Bahnen und Satellitenuhr-Korrekturen für alle GNSS-Satelliten zu rechnen und frei zur Verfügung zu stellen. Leider war diese positive Entwicklung, bedingt durch die wirtschaftlichen Probleme in der russischen Föderation, auch begleitet von vielen Ausfällen in der GLONASS-Konstellation: In den ersten Jahren des neuen Jahrhunderts sank die Zahl der aktiven GLONASS-Satelliten zum Teil unter zehn. Die Jahre 2006–2010 waren dann aber geprägt vom Wiederaufbau der Konstellation, sodass GLONASS 2011 erneut als operationell gemeldet werden konnte. In den Jahren 2000–2015 wurden auch die Entwicklungen für das europäische Galileo und das chinesische Beidou vorangetrieben. Heute stehen damit zwei vollständige und zwei im Aufbau begriffene GNSS zur Verfügung. Das amerikanische GPS wurde für dreidimensionale Positionierung in Echtzeit ausgelegt, mit einer angestrebten Genauigkeit weniger Meter. Diese Zielsetzung wurde von allen weiteren GNSS übernommen. Bild 4 zeigt eine Prinzipskizze eines GPS-Satelliten erster Generation, die *cum grano salis* für alle heute aktiven GNSS-Satelliten Gültigkeit hat. Heutige GNSS-Satelliten bestehen aus einem Zentralkörper (mit Volumina von einigen m^3) und aus Sonnenpaneelen mit Flächen von je etwa $5 \times 2 \text{ m}^2$.

Die Masse der Satelliten beträgt etwa 1000 kg. Die numerischen Werte sind Größenordnungen, die von System zu System um etwa einen Faktor 2 variieren können. Die Sendeantennen sind auf die Erde hin ausgerichtet (z-Achse), die Achse der Sonnenpaneele wird als y-Achse bezeichnet, die zur optimalen Energieerzeugung in der Regel senkrecht zur Richtung Satellit-Sonne stehen sollte. Die Satelliten sind etwa 10 mal größer und 10 mal massiver als jene von Transit. Die korrekte Ausrichtung des Satelliten wird mit Schwungrädern realisiert.

Heute aktive GNSS weisen folgende Eigenschaften auf:

1. Die Satelliten der Systeme senden ständig Signale auf allen Trägerwellen aus, die es dem Empfänger-Oszillator erlauben, den Stand der Satellitenuhr zum Zeitpunkt der Aussendung eines Signals mit dem Stand der Empfängeruhr zum Zeitpunkt des Empfangs desselben Signals zu vergleichen. Die Differenz der Ablesungen dieser beiden Uhren, multipliziert mit der Lichtgeschwindigkeit c , wird Pseudodistanz, englisch *pseudorange* genannt. Die Pseudodistanz ist *die* Observable der Navigation.
2. Die Genauigkeit der Ablesung der Satellitenuhr definiert die Genauigkeit der Positionierung in Echtzeit. Ist der Stand dieser Uhr relativ zur Systemzeit zum Zeitpunkt der Aussendung des Signals mit einer Genauigkeit von $\Delta t = 3 \cdot 10^{-9}$ s bekannt, lässt sich die Distanz Empfänger – Satellit auf $c \cdot \Delta t \approx 1$ m (c ist die Lichtgeschwindigkeit) genau berechnen, falls der Fehler der Empfängeruhr vernachlässigbar ist, was sich wie bei VLBI durch Verwenden von Wasserstoffmasern als Empfängeruhren erreichen ließe – allerdings zu einem horrenden Preis bzw. einer erheblichen Einschränkung in den Anwendungen.
3. Vier oder mehr Satelliten des Systems müssen von jedem Empfänger zu jeder Zeit beobachtet werden können: Damit können die drei Koordinaten des Empfängers und seine Uhrkorrektur bezüglich der Systemzeit bestimmt werden. Mit der Erweiterung der Navigation von einem drei- zu einem vierdimensionalen Problem erhalten wir zusätzlich die Synchronisation der Empfängeruhr auf die Systemzeit des GNSS – und können die Forderung nach Atomuhren in den Empfängern fallen lassen.
4. Zusätzlich zu den Pseudodistanzen werden die aufintegrierten Phasen der Trägerwellen gemessen, dies mit einer Genauigkeit weniger Millimeter (entsprechend einigen Hundertsteln der Träger-Wellenlängen von ca. 20 cm). Die Phasenbeobachtung ist *die* Observable für alle hochgenauen GNSS-Anwendungen.

5. Die Wellenlängen der Trägerwellen werden gegenüber Transit massiv reduziert, vom UKW- in den Mikrowellen-Bereich – um genau zu sein, in den L-Band-Bereich (1–2 GHz).

Ad 1. Die Differenz „Empfängeruhr zum Zeitpunkt des Signalempfangs minus Satellitenuhr zum Zeitpunkt der Signalausendung“ lässt sich als Funktion der Laufzeit des Signals vom Satelliten zum Empfänger und als Funktion der beiden Uhrkorrekturen Δt_R und Δt^S darstellen:

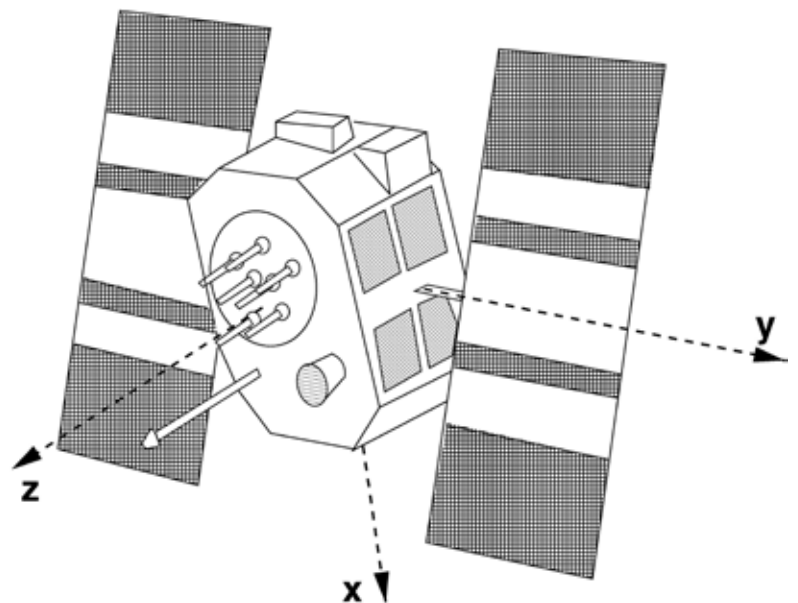
$$t_R - t^S = (t + \Delta t_R) - (t - \Delta t + \Delta t^S) = \Delta t + \Delta t_R - \Delta t^S, \quad (1)$$

wobei t die wahre (Newton'sche oder Einstein'sche) Zeit des Signalempfangs darstellt, Δt die Laufzeit des Signals vom Satelliten zum Empfänger, Δt_R die Empfänger- und Δt^S die Satellitenuhr-Korrektur. Multipliziert man diese Gleichung mit der Lichtgeschwindigkeit c in Vakuum, erhält man die Gleichung für die Pseudodistanz P in Metern:

$$P = c \cdot (t_R - t^S) = \rho + c \cdot (\Delta t_R - \Delta t^S) + \Delta \rho_{\text{ion}} + \Delta \rho_{\text{trop}} + \varepsilon_c, \quad (2)$$

wobei $\rho = c \cdot \Delta t$ die Distanz zwischen dem Satelliten zum Zeitpunkt $(t - \Delta t)$ und Empfänger zum Zeitpunkt t ist sowie $\Delta \rho_{\text{ion}}$, $\Delta \rho_{\text{trop}}$ die Laufzeitverzögerungen des Signals in der Erdatmosphäre, der Ionosphäre respektive Troposphäre darstellen. Die Größe ε_c stellt den Fehler der gemessenen Pseudodistanz dar, wobei der mittlere Fehler für das frei verfügbare Signal bei älteren Empfängern wenige Meter betrug, heute aber, abhängig vom Empfängertyp, oft deutlich kleiner als ein Meter ist.

Bild 4:
Prinzipskizze eines
GNSS-Satelliten
(aus [5], Vol. 1,
Kapitel 3.4)



Zur Navigation, d. h. zur Positionierung in Echtzeit, benötigt man zum Bestimmen der drei Koordinaten und des Synchronisationsfehlers des Empfängers zur Zeit t im Minimum vier Gleichungen vom Typ (2) für vier verschiedene Satelliten.. Die Korrekturen der Satellitenuhren Δt^s und die Satellitenpositionen nimmt man als bekannt an (sie werden vom Satelliten als Broadcast-Message an den Nutzer übermittelt), die ionosphärischen und troposphärischen Laufzeit-Korrekturen der Signale werden durch einfache Modelle erfasst. Die Beobachtungsgleichung für einen Satelliten kann im Längenmaß wie folgt geschrieben werden:

$$c \cdot (t_R - t^s) = \rho + c \cdot \Delta t_R + \varepsilon_c \quad (3)$$

Vier oder mehr Gleichungen vom Typ (3) erlauben es also auch, die Empfängeruhr in Echtzeit auf GNSS-Systemzeit zu synchronisieren – und dies mit der Genauigkeit einiger Nanosekunden (10^{-9} s). Das stellt gegenüber Transit einen Gewinn von etwa einem Faktor 1000 dar.

Ad 2. Die Satellitenuhr-Korrektur muss kleiner als 3 ns sein, damit die Pseudodistanz nicht um mehr als 1 m verfälscht wird. Diese Uhrkorrektur erfolgt durch eine Extrapolation, wobei maximal über das Intervall ΔT zwischen aufeinanderfolgenden Bahn- und Uhrbestimmungen des betreffenden GNSS extrapoliert werden muss. Diese Intervalle ΔT variieren zwischen den verschiedenen GNSS. Rechnet man die Uhrkorrektur in eine Allan-Abweichung um, erhält man einen Wert von $3 \cdot 10^{-9}/43200 \approx 7 \cdot 10^{-14}$ bei einem halben Tag, d. h. bei $\tau = 43200$. Es ist kaum möglich, eine solche Stabilität mit Kristall-Oszillatoren zu erreichen. Aus der etwas willkürlichen aber typischen Genauigkeitsanforderung von 3 ns an die Satellitenuhren folgt zwingend, dass moderne GNSS mit Atomuhren bestückt werden müssen. Die erste Generation GPS- und GLONASS-Satelliten hatten daher kompakte Caesium- oder Rubidium-Atomuhren an Bord. Aktuell verwendet GPS nur noch – allerdings sehr gute – Rubidium-Atomuhren. Für Galileo wurde eine kompakte Version eines Wasserstoffmasers entwickelt, daneben werden auch hier Rubidium-Atomuhren verwendet.

Ad 3. Die gleichzeitige Beobachtbarkeit von vier oder mehr Satelliten muss zu jedem Zeitpunkt und überall auf der Erde gegeben sein. Daraus ergeben sich Anforderungen an die Satelliten-Konstellation: Ein Transit-Satellit sieht aus einer Höhe von 1045 km etwa 7 % der Erdoberfläche. Verlangt man zusätzlich, dass ein Satellit von der Erde aus mindestens unter einer Elevation von 20° beobachtet werden soll, reduziert sich dieser Prozentsatz auf etwa 2 %.

Bei der Bahnhöhe der GPS-Konstellation von 20000 km erhöhen sich diese Prozentzahlen auf 38 % für eine Minimalelevation von 0° respektive

auf 23 % für eine Minimalelevation von 20°. Daher fliegen GNSS-Satelliten heute auf höheren Bahnen. Verlangt man, dass von jedem Punkt der Erde aus mindestens vier Satelliten gleichzeitig sichtbar sein sollen, ergibt sich, dass im MEO-Gürtel (MEO = *Mean Earth Orbit*) mit Umlaufzeiten von etwa einem halben Tag im Minimum etwa 24 Satelliten nötig sind.

Während beim Transit-System polare Bahnen nötig waren, um die ganze Erde abzudecken, hat man bei GNSS im MEO durch die größere Bahnhöhe mehr Freiheitsgrade. So betragen bei GPS, Galileo und Beidou die Bahnneigungen $i \approx 55\text{--}56^\circ$, bei GLONASS $i \approx 65^\circ$, was allerdings bedeutet, dass die maximale Elevation der Satelliten über den Polen für GLONASS 57° beträgt, für alle anderen Systeme etwa 45° . Die Zahl der Bahnebenen und die sich daraus ergebende Zahl der Satelliten pro Bahnebene können im Rahmen des Sinnvollen frei gewählt werden. GPS hat sechs, die übrigen GNSS je drei Bahnebenen.

Ad 4. Bislang haben wir den Eindruck erweckt, dass sich die Observable der aktuellen GNSS grundsätzlich von derjenigen von Transit unterscheidet. Dies ist für Navigationsanwendungen mit Genauigkeitsanforderungen im Meterbereich, wie Navigationshilfen für Fußgänger, Autofahrer, Piloten, etc. auch richtig. Für Aufgaben, bei denen – wie bei VLBI – Millimeter- und/oder Mikrobogensekunden-Genauigkeiten angestrebt werden, ist dies unzureichend. Hier spielen die rekonstruierten Phasen der Trägerwellen die zentrale Rolle. Das Messprinzip ist einfach: Der Empfänger rekonstruiert zu einer vorgegebenen Zeit t seiner Uhr die Phasen der empfangenen Trägerwellen. Die Phasen-Observable φ kann geschrieben werden als:

$$\varphi = \rho + c \cdot (\Delta t_R - \Delta t^s) - \Delta \rho_{\text{ion}} + \Delta \rho_{\text{trop}} + \lambda N + \varepsilon_\varphi, \quad (4)$$

wobei der Term λN die Phasenmehrdeutigkeit (*initial phase ambiguity*) charakterisiert. Der wichtigste Unterschied zur *Code-Observable* (2) besteht darin, dass der mittlere Messfehler von ε_φ lediglich wenige mm beträgt – was einen Genauigkeitsgewinn von mehr als einem Faktor 100 gegenüber den *Code-Messungen* entspricht und den Schritt vom Meter zum Millimeter ermöglicht. Die übrigen Unterschiede findet man beispielsweise in [5, Band 1, Abschnitt 8.5.3] diskutiert. Möchte man die Gleichung (4) für genau einen Zeitpunkt t verwenden, kann die Unbekannte λN nicht von Δt_R getrennt werden. Erst wenn je vier oder mehr Gleichungen vom Typ (4) für verschiedene Zeitpunkte t miteinander kombiniert werden, gelingt diese Trennung, da sich der Term λN erst ändert, wenn der Empfänger den Satelliten nicht mehr verfolgen kann. Die Gleichungen vom Typ (4) sind also wesentlich aufwendiger zu lösen, muss dies doch „*en bloc*“ in einer

Kombination über alle Zeitpunkte geschehen. Der Lohn für die Mühe besteht in wesentlich kleineren Fehlern der zu bestimmenden Parameter.

Anmerkung zur Geschichte: Für wissenschaftliche Anwender waren die mit dem Macrometer-1000 [6] in den frühen 1980er-Jahren auf kurzen Basislinien erzielten Resultate von wegweisender Bedeutung: Mit dem Einfrequenz-Empfänger war es möglich, mit sogenannten Sessionen (Mess-Kampagne) von 1–2 Stunden in Gebieten von 20 km Durchmesser Koordinaten mit einer relativen Genauigkeit von etwa 10^{-6} zu erzielen. So konnte eine Basislinie von 1 km Länge mit einer Genauigkeit von 1–2 mm, eine von 20 km Länge mit einer Genauigkeit von etwa 2 cm vermessen werden. Resultate dieser Qualität waren insbesondere für die Geodäsie und für Vermessungsämter höchst interessant.

Das Macrometer kam der VLBI-Technik sehr nahe: Gemessen und analysiert wurden nur die Phasen, die *Code-Messungen* und die Satelliteninformation wurden ignoriert. Gegenüber VLBI hatte das Macrometer den Vorteil, dass die Signale nicht zufällig, sondern bekannt (zirkular polarisierte Sinus-Wellen bekannter Wellenlänge) und viel stärker waren, was die Verwendung relativ handlicher omnidirektionaler Antennen ermöglichte. Die Genauigkeit der Beobachtungen war bei VLBI und Macrometer etwa vergleichbar, im Bereich weniger Millimeter.

Der erste „vollständige“ für zivile Anwendungen zur Verfügung stehende GPS-Empfänger war der TI-4100 (TI=Texas Instruments). Es stand die ganze Palette der in Echtzeit von den Satelliten ausgesandten Informationen auf beiden Trägerwellen zur Verfügung. Damit wurde es möglich, auch lange Basislinien mit einer hohen relativen Genauigkeit von 10^{-6} oder besser zu vermessen. Auch wurde es möglich, die Bahnen der Satelliten zu verbessern. Im Artikel [7] ist die Auswertung einer vom amerikanischen *National Geodetic Survey* durchgeführten Messkampagne mit Basislinien von Hunderten von Kilometern beschrieben.

Ad 5. Der Übergang vom (400, 150) MHz-Bereich in den (1.57542, 1.22760) GHz-Bereich reduziert die ionosphärische Refraktion um Faktoren von $((1.57542/0.4)^2, (1.22760/0.15)^2) \approx (15.5, 67.0)$, da die ionosphärische Refraktion umgekehrt proportional dem Quadrat der Frequenz ist. Dank dieser Reduktion konnten geodätisch interessante Einfrequenz-Resultate z. B. mit dem Macrometer erzielt werden.

Der IGS: Möglichkeiten und Grenzen der Positionierung mit GNSS

Gegen Ende der 1980er-Jahre wurden für wissenschaftliche Anwendungen immer mehr Netze immer größerer Ausdehnung mit GPS vermessen. Dabei wurde klar, dass die in Echtzeit zur Verfü-

gung stehenden *Broadcast-Bahnen* **der** limitierende Faktor für die Genauigkeit der Auswertungen waren. Diese Einsicht führte zur Gründung des *International GPS Service* (IGS), später umbenannt in *International GNSS Service* – unter Beibehaltung des gleichen Akronyms. Der IGS wurde ab 1989 geplant. 1992 wurde eine drei Monate dauernde Testkampagne durchgeführt, die so vielversprechend war, dass unmittelbar anschließend ein Pilot-Dienst eingerichtet wurde, sodass die IGS-Aktivitäten seit Juni 1992 nie unterbrochen wurden. Im Jahr 1994 wurde der IGS [8, 9] ein offizieller wissenschaftlicher Dienst der Internationalen Assoziation für Geodäsie (IAG).

1989 hatte man nur das Ziel, GPS-Bahnen und Satellitenuhr-Korrekturen relativ zur GPS-Systemzeit zu bestimmen, dies mit einem permanenten Netz von damals etwa dreißig Bodenstationen. Die Stationskoordinaten und ihre Bewegung in einem globalen Referenzsystem sowie die Erdrotationsparameter (ERP) sollten wenn immer möglich vom IERS (*International Earth Rotation and Reference Systems Service*) übernommen werden. Dessen Analysen der Erdrotation beruhten damals hauptsächlich auf VLBI-Beobachtungen.

Die IGS-Pilotphase hatte gezeigt, dass das ursprüngliche Konzept beträchtlich erweitert werden musste: Zusätzlich zu den Satellitenbahnen und -uhren mussten täglich die Koordinaten der Stationen des IGS-Netzes und die der Satellitengeodäsie zugänglichen Erdrotationsparameter bestimmt werden.

Die Erdrotationsparameter legen die Rotation des festen Erdkörpers in einem inertialen Bezugssystem fest. Der Satellitengeodäsie einfach zugänglich ist die Polschwankung, die Wanderung des Rotationspols der Erde auf der Erdoberfläche und die Variationen der Periode der Erdrotation – relativ zu einer streng konstanten Tageslänge. Die übrigen Parameter legen die Lage der Rotationsachse im inertialen Raum fest und sind mit den Methoden der Satellitengeodäsie nur schwer zu bestimmen.

Bild 5a zeigt das IGS-Netz der heute zur Verfügung stehenden mehr als 400 IGS-Bodenstationen, Bild 5b das aus den IGS-Resultaten durch den IERS daraus abgeleitete Geschwindigkeitsnetz. Die Genauigkeit des Bodennetzes steht dem von VLBI in Genauigkeit nicht nach, das IGS-Netz ist aber bedeutend dichter. Die Bilder 5a und 5b dokumentieren auch eindrucksvoll, dass eine enge Zusammenarbeit zwischen dem IGS und dem IERS unbedingt erforderlich ist.

Bild 5c zeigt die vom CODE-Rechenzentrum (CODE = *Center for Orbit Determination in Europe*) bestimmte Polschwankung 2000–2014 mit einer zeitlichen Auflösung von einem Tag in Einheiten von 0,001". Die Genauigkeit beträgt einige zehn μ as. Zum Vergleich findet man in

Bild 5d die mit astrometrischen Methoden im Zeitintervall 1900–1914 bestimmte Polschwankung, wobei die zeitliche Auflösung nur etwa 18 Tage beträgt (C01 Pol von <https://www.iers.org/IERSEN/DataProducts/EarthOrientationData/eop.html>). Die Polkoordinaten konnten damals lediglich auf einige Hundertstel Bogensekunden genau bestimmt werden. Die Bilder 5e und 5f zeigen die Variationen der Tageslängen in den Zeitintervallen 2000–2014 respektive im Jahr 2010

in Millisekunden, wie sie vom CODE-Rechenzentrum bestimmt wurden. Die schwarzen Kurven zeigen die täglichen Schätzwerte, bei den roten Kurven wurden die bekannten zonalen Gezeitenvariationen subtrahiert. Die roten Kurven können zum größten Teil durch Drehimpuls-Austausch zwischen fester Erde und Atmosphäre erklärt werden (siehe dazu z. B. [5], Vol. 2, Bild 2.47). Sie zeigen aber auch Unzulänglichkeiten der Gezeitenmodelle.



Bild 5a: Das permanente IGS-Netz 2016

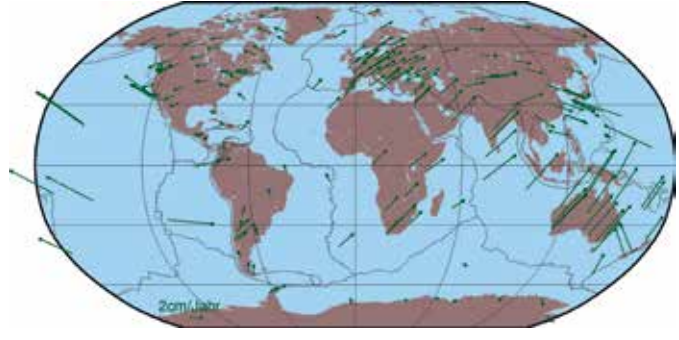


Bild 5b: Geschwindigkeitsfeld des IGS14-Feldes, aus IGS-Beiträgen, berechnet vom IERS (IGSMail 7399)

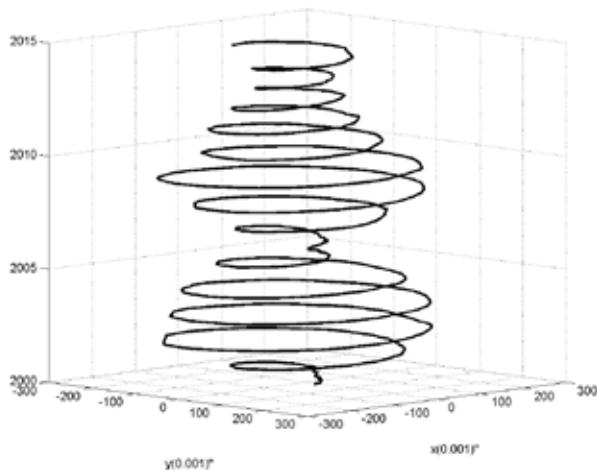


Bild 5c: Polschwankung 2000–2014

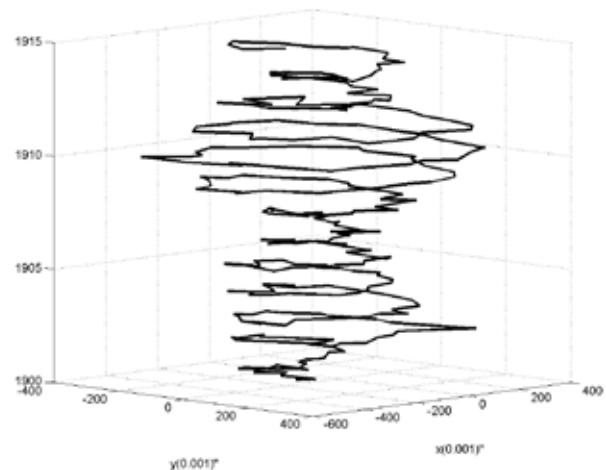


Bild 5d: Polschwankung 1900–1914

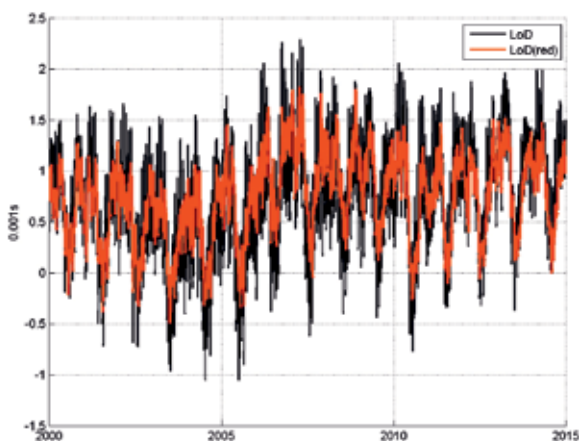


Bild 5e: Tageslängenvariation 2000–2014



Bild 5f: Tageslängenvariation 2010

Parameter	Mittlere Fehler	Bemerkungen
Broadcast-Bahnen (GPS)	1 m	Werte früherer Jahre deutlich schlechter
IGS Final Orbits (GPS, GLO-NASS)	2,5 cm (3 cm)	Werte für GPS (GLONASS)
Broadcast-Satellitenuhren (GPS)	2,5 ns – 5 ns	mit S/A mehr als 50–100 mal größer
IGS-Final-Satellitenuhren	75 ps	$1 \text{ ps} = 10^{-12} \text{ s}$, $75 \text{ ps} \cdot c \approx 2,2 \text{ cm}$
Polkoordinaten x und y	30 μs	etwa 1mm auf der Erdoberfläche
Tageslängenexzess	10 μs	etwa 4,6 mm/Tag
IGS-Stationskoordinaten (final)	3 mm (6 mm)	horizontale (vertikale) Koordinaten

Tabelle 2:
IGS-Produkte und
ihre geschätzten
mittleren Fehler,
Stand 2017 ([http://
www.igs.org/pro-
ducts](http://www.igs.org/products))

Tabelle 2 fasst die Qualität einiger der wichtigsten IGS-Produkte zusammen, genauer gesagt der sogenannten endgültigen oder *final-Produkte*, die etwa zwei Wochen im Nachhinein zur Verfügung stehen. Als Vergleich wurden auch die *Broadcast-Produkte* von GPS aufgenommen. Die Tabelle zeigt, dass die Qualität verschiedener Produkte ähnlich ist, wenn man sie in vergleichbare Einheiten umwandelt. So entsprechen 30 μs auf der Erdoberfläche etwa einem Millimeter, eine Satellitenuhr-Korrektur von 75 ps induziert eine Abweichung von etwa 2,2 cm in der Distanz zwischen Satellit und Station.

Die Genauigkeit, mit der die Satellitenuhr-Korrektur bestimmt werden kann, entspricht auch der Genauigkeit, mit der die Uhrenkorrektur der Empfänger bestimmt werden können, wenn dies mit wissenschaftlichen Programmsystemen geschieht.

Die Entwicklung des IGS zu einem interdisziplinären wissenschaftlichen Dienst findet man in Referenz [8] beschrieben, seinen heutigen Stand in [9]. Die Multi-GNSS-Entwicklungen im IGS sind in [10] behandelt.

GNSS und Atomuhren: Zusammenfassung

Ist Navigation ohne Atomuhren auf den GNSS-Satelliten überhaupt möglich? Da das Transit-System auf Kristall-Oszillatoren beruhte, ist man geneigt, die Frage zu bejahen. Dem Argument ist aber zu widersprechen, da Transit eine Navigation im heutigen Sinn nicht erlaubte.

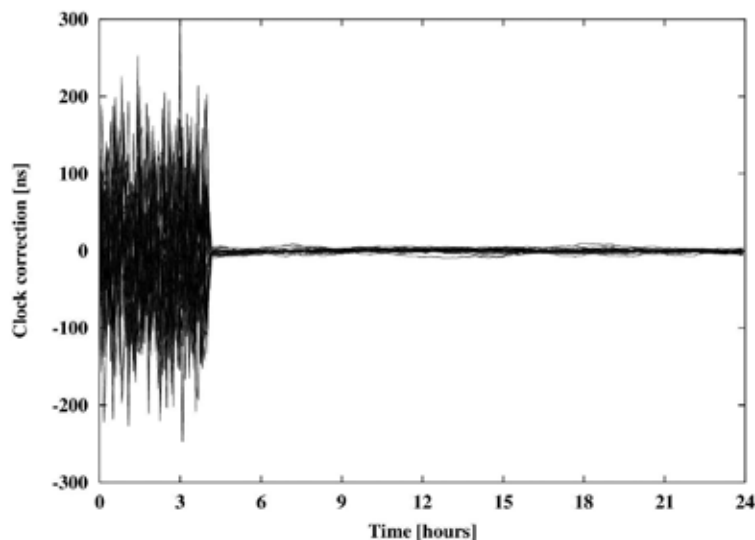
Moderne Satelliten-Navigation ist äquivalent zu momentaner dreidimensionaler Positionierung, dies mit „billigen“ Empfängern. Billige Empfänger bedingen günstige Uhren, womit nach heutigen Verhältnissen Atomuhren in Empfängern auszuschließen sind. Dreidimensionale Positionierung mit Mikrowellensystemen verlangt die gleichzeitige Beobachtung von mindestens vier Satelliten. Damit in Echtzeit die Genauigkeitslimite von wenigen Metern erreicht werden kann, erfordert dies die Präzisierbarkeit aller Satellitenuhr-Korrekturen über Zeitintervalle von mehreren Stunden mit einer Genauigkeit von etwa 3 ns, eine Anforderung, der nur Atomuhren gerecht werden können.

Moderne Satelliten-Navigation bedeutet heute *auch* Navigieren mit mehreren Satellitensystemen. Nach Tabelle 1 stehen zwei vollständige Systeme, GPS und GLONASS, und zwei im Aufbau begriffene, Beidou und Galileo, zur Verfügung. Eine regionale Variante von Beidou ist schon operationell, mit 18 aktiven Galileo-Satelliten (Stand Juni 2017) ist mehr als die Hälfte der Galileo-Konstellation (27 Satelliten) bereits heute verfügbar.

Gerade für die Navigation ist der Multi-GNSS-Aspekt sehr bedeutsam, steigt doch die Genauigkeit und vor allem die Zuverlässigkeit mit der Zahl gleichzeitig beobachteter Satelliten. Die Kombination verschiedener GNSS in der Auswertung stellt aber eine nicht geringe Herausforderung dar, wie die Referenz [10] eindrücklich belegt. Der Aspekt Multi-GNSS wäre hochinteressant, geht aber über unsere, sich auf prinzipielle Aspekte beschränkende Betrachtungen hinaus.

Dass Satelliten-Navigation ohne Atomuhren kaum möglich ist, wird durch Bild 6 (Abbildung 2.7 aus [11]) belegt, das für den 2. Mai 2000 die Satellitenuhr-Korrekturen aller damals aktiven GPS-Satelliten als Funktion der Tageszeit zeigt. Am 2. Mai 2000 wurde S/A für das ganze GPS abgeschaltet. Offenbar fand dieses für GPS denkwürdige Ereignis etwa um 4^h UTC statt.

Bild 6:
Korrekturen der
Broadcast-Uhren
relativ zu den
geschätzten Uhr-
korrekturen des
Code-Rechenzen-
trums am 9. Mai
2000, als S/A auf
GPS ausgeschaltet
wurde (Abbildung
2.7 aus [11])



Vor diesem Zeitpunkt waren die Satellitenuhren dem Navigator im Mittel nur auf etwa 70 ns genau bekannt, was einem Fehler in der Distanz Empfänger – Satellit von $70 \cdot c \approx 20$ m entspricht und damit ein Navigieren mit einer Genauigkeit von wenigen Metern verhinderte. Das Zeitintervall von 0^h–4^h UTC mit S/A charakterisiert eine Satelliten-Navigation mit „sehr schlechten“ Satellitenuhren. Danach erhöhte sich die Genauigkeit der Uhrprädiktion auf wenige ns.

S/A hat die Aktivitäten des IGS von 1992 bis zum 8. Mai 2000 praktisch nicht beeinflusst, da diese auf Phasenbeobachtungen und auf der Bestimmung aller für Positionierung und Navigation relevanten Parameter beruht, einschließlich der Satellitenuhr-Korrekturen.

Es seien zum Abschluss noch zwei Aspekte erwähnt, bei denen Atomuhren unbedingt erforderlich sind:

1. Jedes GNSS muss eine systemeigene Zeitskala realisieren, die eindeutig auf UTC (und damit auf Atomzeit) rückführbar sein muss. Diese Rückführung wird durch Atomuhren auf der Erde realisiert, die einerseits an die internationale Atomzeit angebunden, andererseits mit dem betreffenden GNSS verknüpft sind.
2. Das sogenannte *Precise Point Positioning* (PPP) hat sich für viele GNSS-Anwendungen, beispielsweise auch für die genaue Bestimmung von Bahnen von *Low Earth Orbiters* (LEOs), als extrem wichtige und effiziente Methode etabliert. Die Trajektorie des GNSS-Empfängers wird bestimmt mithilfe seiner Messungen – unter Verwendung der präzisen GNSS-Satellitenbahnen- und -uhr-Korrekturen von IGS oder von einem seiner Rechenzentren. Nun werden die GNSS-Satellitenuhr-Korrekturen im Allgemeinen nur in einem Raster von 30 Sekunden bis 5 Minuten zur Verfügung gestellt, was bedeutet, dass die Stabilität der Uhren in diesen Zeitintervallen so hoch sein muss, dass eine Interpolation innerhalb dieser Intervalle ohne Genauigkeitsverlust möglich ist. Beispielsweise verlangt dies nach einer relativen Frequenzschwankung von 10^{-13} oder weniger bei einer Mittelungszeit von $\tau \approx 30$ s, was dann einem Fehler von etwa 1 mm in der Distanz Empfänger – Satellit entspricht.

Schließlich sei darauf hingewiesen, dass die meisten Auswerte-Algorithmen das Potenzial der heutigen, hochpräzisen Satellitenuhren nicht optimal ausschöpfen, da jede Satellitenuhr-Korrektur als Unbekannte eingeführt und geschätzt wird, anstatt dass zusätzlich die Parameter von Satellitenuhr-Modellen geschätzt und die Variationen

der Epochen- und Satelliten-spezifischen Uhrenkorrekturen relativ zu diesem Modell beschränkt werden – dies unter Verwendung bekannter stochastischer Modelle für die Atomuhren. Durch das skizzierte Verfahren würden die mittleren Fehler der verbleibenden Parameter gegenüber den heute üblichen Verfahren beträchtlich reduziert.

Dank

Herrn Rolf Dach, dem Leiter der Forschungsgruppe Satellitengeodäsie am Astronomischen Institut der Universität Bern, danke ich für die sehr gründliche Durchsicht des Manuskriptes und für viele wertvolle Anregungen. Herrn Oliver Montenbruck, dem Leiter des IGS-Multi-GNSS-Experimentes und dem Herausgeber des soeben erschienenen Handbuches über GNSS [12], danke ich für seine kritische Durchsicht des Manuskriptes. Herrn Heinz Spahni, meinem alten Kollegen und Freund, danke ich für die Durchsicht des Manuskriptes aus der Sicht des Physikers.

Referenzen

- [1] Parkinson, B.W., Powers S.T. (2010) Part 1, „*The origins of GPS, and the pioneers who launched the system*“, GPS World, May 1, 2010
- [2] Parkinson, B.W., Powers S.T. (2010) Part 2, „*The origins of GPS, fighting to survive*“, GPS World, June 1, 2010
- [3] Proceedings of the sixth international symposium of satellite positioning, 17–20 März 1992, Ohio State University.
- [4] J. Kouba, „*Geodetic Doppler positioning and application to Canadian test adjustments*“, Phil. Trans. Royal Soc. Lond. A **294** (1980), pp. 271–276
- [5] G. Beutler, „*Methods of Celestial Mechanics*“, Vol. I and II, Springer Verlag, Heidelberg, Berlin, New York, (2005)
- [6] Y. Bock, R.I. Abbot, C.C. Counselman, S.A. Gourevitch, R.W. King, A.R. Paradis, „*Geodetic Accuracy of the Macrometer Model V-1000*“, Bull. Géod. **58** (1984), pp. 211–221
- [7] G. Beutler, I. Bauersima, W. Gurtner, M. Rothacher, T. Schildknecht, „*Evaluation of the Alaska Global Positioning System Campaign with the Bernese Software*“, J. Geophys. Res., **92**, B2 (1987), pp. 1295–1303
- [8] G. Beutler, M. Rothacher, S. Schaer, T.A. Springer, J. Kouba, R.E. Neilan, „*The International GPS Service (IGS): An interdisciplinary service in support of Earth sciences*“, Adv. Space Res. **23** (1999), pp. 631–653
- [9] J.M. Dow, R.E. Neilan, C. Rizos, „*The International GNSS Service in a changing landscape of Global Navigation Satellite Systems*“, J. Geod. **83** (2009), pp. 191–198

- [10] O. Montenbruck, P. Steigenberger, L. Prange, Z. Deng, Q. Zhao, F. Perosanz, I. Romero, C. Noll, A. Stürze, G. Weber, R. Schmid, K. MacLeod, S. Schaer, „*The Multi-GNSS Experiment (MGEX) of the International GNSS Service (IGS) – Achievements, Prospects and Challenges*“, Adv. Space Res., **59**(7) (2017), pp. 1671–1697
- [11] R. Dach, S. Lutz, P. Walser, P. Fridez (eds), „*The Bernese GPS software version 5.2*“, Astronomical Institute, University of Bern, (2015), doi: 10.7892/boris.72297
- [12] P.J.G. Teunissen, O. Montenbruck (eds), „*Springer Handbook of Global Navigation Satellite Systems*“, Springer Verlag Heidelberg, Berlin, (2017), New York, doi 10.1007/978-3-319-42928-1

Ultrastabiles Differenz-Interferometer Serie SP-DI

- Interferometer für Differenzmessungen höchster Stabilität mit einer Auflösung von 20 pm
- Temperaturkoeffizient: < 20 nm/K
- Längenmessbereich: 2 m, Winkelmessbereich: $\pm 1,5'$
- Für Materialuntersuchungen und Langzeitmessungen
- Auch mehrkanalige Ausführung möglich



SIOS Meßtechnik GmbH
Am Vogelherd 46
D-98693 Ilmenau
Tel.: +49-(0)3677/6447-0
E-mail: contact@sios.de
www.sios.de



Einstein und die Zeit

Claus Kiefer*

1. Zeit vor Einstein

Wie kaum ein anderer Name ist der von Albert Einstein (1879 bis 1955) mit dem modernen Weltbild der Physik verbunden. Seine Spezielle und seine Allgemeine Relativitätstheorie bilden zusammen mit der Quantentheorie die Grundlage unseres physikalischen Weltverständnisses. Unsere Vorstellungen von Raum und Zeit haben durch seine Arbeiten eine wesentliche Entwicklung erfahren, was auf die Zeit vielleicht noch mehr zutrifft als auf den Raum.

Was ist das Wesen des Einstein'schen Zeitbegriffs? Und ist zu erwarten, dass dieser eine weitere Wandlung erfahren wird? Unsere Auffassung von dem, was wir Zeit nennen, hat sich im Laufe der Jahrhunderte stetig verändert, beeinflusst natürlich durch Entwicklungen im Denken allgemein und in den Wissenschaften im Besonderen. Diese Entwicklungen sind noch nicht an ihr Ende gelangt.

In der Antike war Zeit untrennbar mit Bewegung verknüpft – war quasi die Zahl der Bewegung. Das Maß hierfür bildete die Bewegung der Himmelskörper in der Form der täglichen Umdrehung von Sonne und Fixsternsphäre und der jährlichen Sonnenbewegung. Diese Bewegungen waren periodisch und abzählbar – etwas anderes war zu jener Zeit nicht vorstellbar, insbesondere kein Kontinuum. Das Gleichmaß der Bewegungen musste für diesen Zweck absolut vorgegeben sein, bei Aristoteles etwa in Gestalt des „unbewegten Bewegers“ als deren Garant. Laut Hans Blumenberg ist es für Aristoteles „der Begriff der Zeit, der die Existenz der absolut stetigen Bewegung und ihres Garanten impliziert.“ [1]

Das Weltbild des Aristoteles hielt fast zwei Jahrtausende. Der Umschwung setzte mit den revolutionären Ideen des Kopernikus ein, der die Erde dem Mittelpunkt der Welt entrückte und sie zu einem Himmelskörper unter anderen Himmelskörpern machte. Vor Kopernikus war die himmlische Welt strikt von der irdischen getrennt – der Astronomie ging es nur um Beschreibung und um konsistente Hypothesen; sie besaß keinen

Wahrheitsanspruch im Sinne der auf das Irdische beschränkten Physik. Das spiegelt sich noch in der anonymen Vorrede des Andreas Osiander zu Kopernikus' Hauptwerk wider und dessen unautorisierter Abänderung des von Kopernikus beabsichtigten Titels „Über die Kreisbewegungen der Weltkörper“ (*De revolutionibus orbium mundi*) in „Über die Kreisbewegungen der Himmelskörper“ (*De revolutionibus orbium coelestium*). Nur bei dem Wort Himmelskörper könnten die alten astronomischen Hypothesen beibehalten werden; Weltkörper hingegen bringt Erde, Planeten und Sterne unter einen Hut. Kopernikus hatte jedoch sehr wohl einen Wahrheitsanspruch – für ihn war die astronomische Welt mehr als eine konsistente Sammlung von Hypothesen. In Blumenbergs Worten: „Kopernikus stellte heraus, dass eine einheitliche Theorie der Welt in durchgehender Rationalität möglich sei.“ [1] Es war genau dieser Wahrheitsanspruch, der den Nachfolgern des Kopernikus Probleme mit der kirchlichen Autorität bereiten sollte. Newtons bekannte Sentenz „*Hypotheses non fingo*“ („Ich stelle keine Hypothesen auf“) ist in diesem Zusammenhang zu verstehen.

Johannes Kepler teilte diesen Wahrheitsanspruch. Die neue Stellung der Sonne in der Weltmitte hatte für ihn aber nur dann einen Sinn, wenn von ihr ein realer Einfluss auf die Himmelskörper ausging. Kepler musste erkennen, dass in einer Astronomie ohne Hypothesen die Kreisbewegungen der Planeten nicht mehr zu halten waren – diese bewegen sich stattdessen auf Ellipsen, in deren einem Brennpunkt die Sonne steht; die Bewegung des Mars' ließ daran keinen Zweifel. Der als real erkannte Einfluss der Sonne führte bei Isaac Newton dann zu dem Gesetz der universellen Gravitation.

Die Kopernikanische Wende hatte gewichtige Auswirkungen auf den Zeitbegriff. Wenn die Himmelsbewegungen die Zeitgeber waren, sollte dann nicht die Zeit stillstehen, nachdem diese wegen der realen Bewegung der Erde als Scheinbewegungen entlarvt worden waren? Sollte dann nicht auch die Bewegung einer Töpfer-

* Prof. Dr. Claus Kiefer, Institut für Theoretische Physik, Universität zu Köln, Zülpicher Straße 77, 50937 Köln

scheibe unmöglich werden? Dieser Schluss konnte vermieden werden, wenn sich der Zeitbegriff von diesen Bewegungen löste, was bei Newton auf den Begriff der absoluten Zeit führte. Diese „verfließt an sich und gleichförmig und ohne Bezug auf irgend einen äußeren Gegenstand“. Trotz berechtigter Kritik (etwa durch Leibniz) war es dieser Zeitbegriff, der bis zu Einsteins Arbeiten Bestand haben sollte. Ohne ihn waren die Formulierung der Newton'schen Gesetze und deren erfolgreiche Anwendung auf die Himmelsmechanik undenkbar.

Die Ablösung der Zeit von den Himmelsbewegungen hatte sich freilich schon früher angekündigt, und zwar mit dem Aufkommen der ersten Räderuhren um 1300. Das Zifferblatt und die periodische Bewegung der Zeiger spiegelte zwar die Tagesbewegung der Fixsternsphäre und der Sonne wider, doch wurde die Bewegung durch einen rein mechanischen Mechanismus verursacht. Diese Bewegung wurde später durch Newtons Gesetze beschrieben, in welche die absolute Zeit t explizit und unveränderlich eingeht. Es ist vielleicht kein Zufall, dass um diese Zeit erste Überlegungen zur Möglichkeit eines Perpetuum mobile zu finden sind, was die Konstruktion einer ewig laufenden Uhr ermöglichen würde. (An den Energiesatz war zu jener Zeit noch nicht zu denken.)

Die historische Rolle der Himmelsbewegungen als Zeitgeber hat trotz dieser Entwicklungen in der Astronomie noch lange eine Rolle gespielt, und zwar in Gestalt der mittleren Sonnensekunde und später der Ephemeriden-Sekunde. Erstere beruht auf der Tagesdrehung der Erde, letztere auf deren Jahresbewegung um die Sonne. Abgelöst wurden diese Einheiten von der 1967 eingeführten Atomsekunde, die auf einem bestimmten Übergang atomarer Niveaus des Cäsiumatoms beruht und deren Jubiläum wir in diesem Heft gedenken. Um Atomzeit mit Ephemeriden-Zeit in Einklang zu bringen, werden seit 1972 Schaltsekunden eingeführt.

Zu dem ersten Ereignis dieser Art merkt Blumenberg lakonisch an: „Verglichen mit den Kalenderstreitigkeiten früherer Jahrhunderte und ihrer Wirkung auf das Bewußtsein vom Ordnungszustand der Welt ließ sich bei diesem Ereignis keine öffentliche Betroffenheit ausmachen.“ [1]

2. Einstein I: Spezielle Relativitätstheorie

In der absoluten Raum-Zeit Newtons lassen sich Bewegungen beliebiger Art beschreiben. Dennoch ist eine bestimmte Klasse ausgezeichnet: die Klasse gleichförmig geradliniger Bewegungen, die kräftefreie Körper charakterisieren. Schon Galilei bemerkte, dass diese physikalisch nicht unterscheidbar sind, es also keinen Zustand absoluter

Ruhe geben kann. Klar ausgesprochen hat diesen Sachverhalt zum ersten Mal Ludwig Lange 1885, der für die durch diese Bewegungen ausgezeichneten Bezugssysteme den heute gebräuchlichen Namen Inertialsysteme prägte (vom lateinischen „inertia“ für Trägheit).

Deren ausgezeichnete Natur ist Inhalt des Trägheitsgesetzes, in der Formulierung von Domenico Giulini: „Ein kräftefreier Körper verharrt im Zustand der Ruhe oder der gleichförmig geradlinigen Bewegung, sofern man das räumliche Referenzsystem und die Uhr geeignet wählt.“ [2] Der Nachsatz ist wichtig, da für allgemeine Bezugssysteme (z. B. für rotierende Systeme) kräftefreie Bewegungen nicht geradlinig sind.

Die Gleichwertigkeit der Inertialsysteme nennt man Relativitätsprinzip. Dessen allgemeine Gültigkeit schien bis zur Aufstellung der Theorie des Elektromagnetismus durch den schottischen Physiker James Clerk Maxwell um 1865 unbestritten zu sein. Maxwells Theorie jedoch sagte die Existenz elektromagnetischer Wellen voraus, die sich im Vakuum mit der Lichtgeschwindigkeit c ausbreiten. Wellenbewegungen konnte man sich im 19. Jahrhundert nur in einem Medium denken, was zur Vorstellung des sogenannten Äthers führte, der unsichtbar, aber von enormer Festigkeit das Universum durchdringen sollte. Maxwells Gleichungen würden dann nur in einem ausgezeichneten System gelten, einem System, in dem der Äther ruht; das Relativitätsprinzip wäre dann verletzt.

Diese Situation bereitete Ende des 19. Jahrhunderts großes Unbehagen. Sollte man das Relativitätsprinzip aufgeben oder sollte man die Maxwell'sche Theorie abändern? In einem genialen Schachzug konnte der junge Angestellte am Berner Patentamt Albert Einstein 1905 zeigen, dass Relativitätsprinzip und Maxwelltheorie beide gelten, wenn man ein weiteres Postulat hinzufügt: das Postulat der Konstanz der Lichtgeschwindigkeit, das heißt, der Unabhängigkeit der Lichtgeschwindigkeit vom Bewegungszustand von Lichtquelle und Beobachter. In seiner Arbeit schreibt Einstein: „Die Einführung eines „Lichtäthers“ wird sich insofern als überflüssig erweisen, als nach der zu entwickelnden Auffassung weder ein mit besonderen Eigenschaften ausgestatteter „absolut ruhender Raum“ eingeführt, noch einem Punkte des leeren Raumes, in welchem elektromagnetische Prozesse stattfinden, ein Geschwindigkeitsvektor zugeordnet wird.“ [6]

Die Endlichkeit der Lichtgeschwindigkeit c war bereits im 17. Jahrhundert festgestellt worden. Dies gelang dem dänischen Astronomen Ole Rømer

1676 durch Beobachtungen des Jupitermondes Io unter Ausnutzung der von Cassinis wenige Jahre zuvor veröffentlichten Ephemeriden für die Jupitermonde.

Er konnte bereits einen relativ genauen Wert für c angeben. Seit 1983 legt man die Lichtgeschwindigkeit auf den exakten Wert $c = 299.792.458$ m/s fest und definiert hiermit die Längeneinheit Meter als diejenige Strecke, die das Licht im Vakuum in $1/299.792.458$ Sekunde zurücklegt. (Die Sekunde wird wie erwähnt durch die Niveauübergänge im Caesiumatom definiert.) Diese Festlegung ist nur möglich, weil eben c eine universelle Konstante ist.

Einstein löste den scheinbaren Konflikt zwischen Relativitätsprinzip und Elektrodynamik also nicht durch eine komplizierte Umformulierung der Elektrodynamik, sondern durch eine scharfsinnige Kritik an den Grundbegriffen Raum und Zeit. In Newtons Vorstellung von Raum und Zeit war die Frage, ob zwei Ereignisse an unterschiedlichen Orten gleichzeitig stattfinden oder nicht, absolut entscheidbar und unabhängig von jedem Bezugssystem. Bei Existenz einer absoluten Geschwindigkeit sind Raum und Zeit jedoch nicht mehr unabhängig, und die Gleichzeitigkeit zweier Ereignisse hängt vom Bewegungszustand des Beobachters ab. Dieser Sachverhalt lässt sich am Beispiel der sogenannten Lichtuhr veranschaulichen (Bild 1).

In der Mitte zwischen zwei parallelen Spiegeln befindet sich eine Lichtquelle, die zu einem bestimmten Zeitpunkt Lichtblitze in entgegengesetzte Richtung aussendet. Für einen Beobachter, der in bezug auf die Spiegel ruht, treffen die Blitze gleichzeitig bei den Spiegeln ein, werden dort reflektiert und treffen wieder gleichzeitig bei der Quelle ein. Nicht so für einen relativ dazu bewegten Beobachter. Für diesen nähert sich der linke Spiegel dem Blitz, während sich der rechte davon entfernt. Da die Lichtgeschwindigkeit jedesmal dieselbe ist, trifft der linke Blitz *vor* dem rechten auf dem jeweiligen Spiegel ein. (In der Newtonschen Sichtweise würde sich der linke Blitz so verlangsamen und der rechte so viel schneller werden, dass die Gleichzeitigkeit des Eintreffens bewahrt bliebe.) In der Mitte treffen beide wieder gleichzeitig ein - die Gleichzeitigkeit am gleichen Ort gilt weiter.

In der Praxis definiert man Gleichzeitigkeit dadurch, dass die mit den Ereignissen gleichzeitigen „Zeigerstellungen“ von Uhren übereinstimmen müssen, die sich an den Orten der Ereignisse befinden *und die synchronisiert worden sind*. Ganz wesentlich ist hier die Existenz einer Synchronisationsvorschrift für Uhren, die bezüglich eines ausgewählten Bezugssystems erfolgt.

Relativitätsprinzip und Konstanz der Lichtgeschwindigkeit bilden die Grundlage von Einsteins Spezieller Relativitätstheorie (SRT). Hieraus folgen die auf den ersten Blick paradox anmutenden Effekte wie Zeitdilatation und Längenkontraktion. Tatsächlich handelt es sich hier nur um grundlegende Eigenschaften der Raum-Zeit und

deren Konsequenzen für reale Uhren und reale Maßstäbe.

Der Mathematiker Hermann Minkowski hat 1908 eine elegante Formulierung der vierdimensionalen Raum-Zeit (drei Raumdimensionen und eine Zeitdimension) vorgelegt, die Einstein als Ausgangspunkt für seine Allgemeine Relativitätstheorie dienen sollte.

Seit der Aufstellung der SRT wurde diese unzähligen experimentellen Tests unterworfen, die sie alle mit Bravour bestanden hat (siehe etwa die Diskussion in [2]). In Form moderner Positionierungssysteme wie GPS oder Galileo haben die Effekte der SRT Eingang in den Alltag gefunden, da in deren Beschreibung die Konstanz der Lichtgeschwindigkeit direkt eingeht und es bei einer Newtonschen Beschreibung zu Fehlweisen von mehreren Kilometern pro Tag kommen würde (siehe unten).

Die SRT bildet auch die Grundlage aller physikalischen Wechselwirkungen, insbesondere des am Beschleuniger Large Hadron Collider (LHC) in Genf untersuchten Standardmodells der Elementarteilchenphysik - mit einer Ausnahme: der Gravitation. Das führt uns zu Einsteins zweitem großen Wurf: seiner Allgemeinen Relativitätstheorie von 1915.

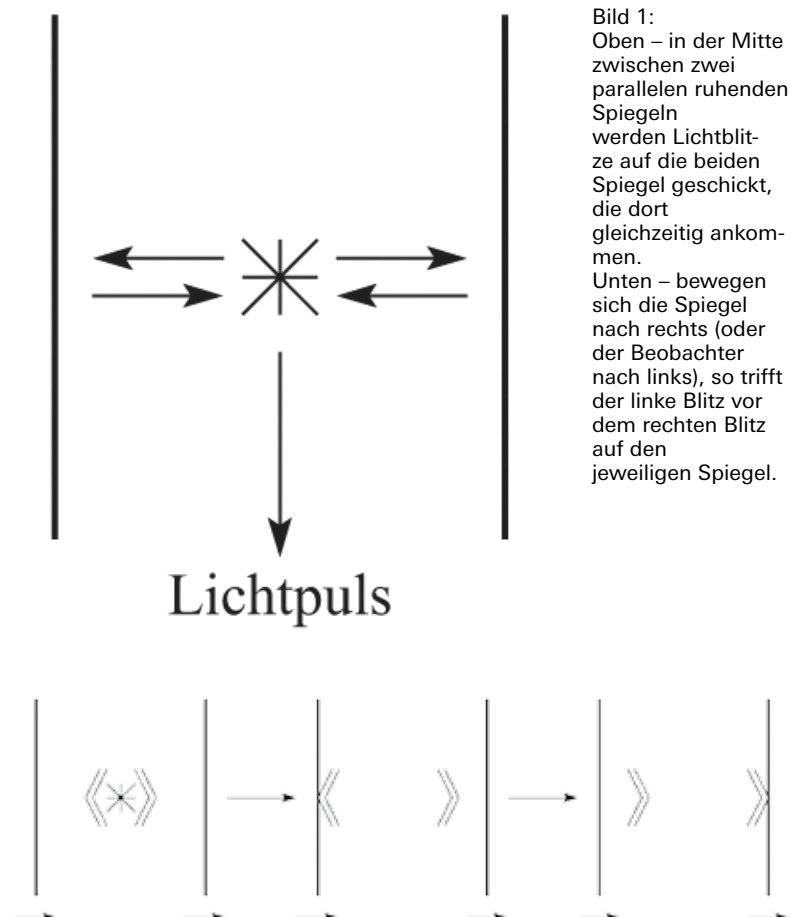


Bild 1:
Oben – in der Mitte zwischen zwei parallelen ruhenden Spiegeln werden Lichtblitze auf die beiden Spiegel geschickt, die dort gleichzeitig ankommen.
Unten – bewegen sich die Spiegel nach rechts (oder der Beobachter nach links), so trifft der linke Blitz vor dem rechten Blitz auf den jeweiligen Spiegel.

3. Einstein II: Allgemeine Relativitätstheorie

Zu Einsteins Zeiten war die Gravitation die einzig bekannte fundamentale Wechselwirkung neben dem Elektromagnetismus. Es war daher für Einstein naheliegend, seine SRT auch auf die Gravitation auszudehnen. Bei seinen Bemühungen musste er freilich einsehen, dass dies nicht möglich ist. Woran liegt das? Bereits Galilei und Newton war bekannt, dass alle Körper im Gravitationsfeld gleich schnell fallen (sofern man den Luftwiderstand unberücksichtigt lässt). Sie haben diese Tatsache aber, in Einsteins Worten, nur registriert, nicht interpretiert. Wenn man in ein frei fallendes Bezugssystem geht, so kann man die Gravitation quasi ausschalten, da alle Körper in diesem System mitfallen und kein Effekt der Gravitation sichtbar ist. Die Situation ist äquivalent zu der in einem schwerkraftfreien Raum, weshalb man hier vom Äquivalenzprinzip spricht. Umgekehrt entspricht die Situation eines etwa auf der Erde stehenden Beobachters einem solchen, der sich mit Erdbeschleunigung in einer Rakete durch den leeren Raum bewegt (für eine ausführliche Diskussion siehe etwa [5]).

Vom Standpunkt des Innenraums der Rakete aus kann die Situation nicht von der Situation einer auf der Erde ruhenden Rakete unterschieden werden. Einstein formulierte das Äquivalenzprinzip dahingehend, dass alle physikalischen Gesetze in einem schwerkraftfreien gleich denen in einem frei fallenden System seien. Zustände freien Falls können heute nicht nur auf Raumstationen, sondern bereits mit normalen Flugzeugen während eines sogenannten Parabelfluges realisiert werden. Bild 2 zeigt den bekannten Physiker Stephen Hawking während eines solchen Fluges.

Ein wichtiger Punkt ist nun, dass solche frei fallenden Systeme nicht zu groß sein dürfen, um das Äquivalenzprinzip noch zu erfüllen. Betrachtet man etwa einen frei fallenden Aufzug auf der

Nordhalbkugel der Erde und einen auf der Südhalbkugel, so sind diese zwar jeweils gravitationsfrei, relativ zueinander aber beschleunigt, da sich beide in Richtung Erdmittelpunkt bewegen. Der Effekt der Gravitation ist also erkennbar, wenn die entsprechenden Systeme groß genug sind.

Solange ein frei fallendes Bezugssystem klein genug ist, um die Gravitation auszuschalten, realisiert es ein lokales Inertialsystem. Das bedeutet, dass in ihm wieder die Gesetze der SRT gelten, kräftefreie Körper sich also auf Geraden bewegen. Aber eben nur lokal! Bei Anwesenheit von Gravitation passen die lokalen Inertialsysteme global nicht mehr zusammen, was nichts anderes heißt, als dass es kein globales Inertialsystem mehr geben kann und die SRT nicht mehr global gültig ist.

Einstein erkannte, dass das Nichtzusammenpassen der lokalen Inertialsysteme bedeutet, dass die vierdimensionale Raum-Zeit nicht mehr flach ist, analog wie das Nichtzusammenpassen etwa von Briefmarken auf einer Kartoffel Ausdruck der Krümmung der Kartoffeloberfläche ist. Es sollte nach dieser Einsicht noch acht Jahre dauern, bis Einstein 1915 seine Feldgleichungen der Gravitation (die berühmten Einstein-Gleichungen) vorlegen konnte. Grund für diese Dauer war die Kompliziertheit der Mathematik gekrümmter vierdimensionaler Raum-Zeiten, die sich Einstein mühsam aneignen musste, nicht zuletzt mithilfe des befreundeten Mathematikers Marcel Grossmann. Die Feldgleichungen sind die zentralen Gleichungen der Allgemeinen Relativitätstheorie (ART), aus denen alle bisher bekannten Phänomene der Gravitation folgen.

Die Einstein-Gleichungen verknüpfen auf elegante Weise die Geometrie der Raum-Zeit mit Energie und Impuls von Materie und nicht-gravitativen Feldern. Agiert Gravitation in der Klassischen Mechanik als geheimnisvolle Kraft, die instantan zwischen beliebig weit entfernten Körpern wirkt, so ist sie in der ART Ausdruck der Geometrie von Raum und Zeit, und zwar der lokalen Geometrie. Es gibt keine Fernwirkung.

In gekrümmten Räumen gibt es Effekte, die man im flachen Raum nicht kennt. Das lässt sich bereits in dem einfachen Beispiel einer zweidimensionalen Kugeloberfläche (statt einer vierdimensionalen Raum-Zeit) erkennen (siehe Bild 3).

Bewegt man einen Pfeil durch Paralleltransport entlang der eingezeichneten Strecke vom Pol zum Äquator und entlang eines anderen Längengrades zurück zum Pol, so stimmt die Pfeilrichtung nicht mehr mit der Ausgangsrichtung überein. Tatsächlich ist die Abweichung ein Maß für die Krümmung der umlaufenen Fläche. In der Raum-Zeit zeigt sich dieser Effekt beispielsweise in der Präzession der Achse eines Kreisel, der in einem Satelliten um die Erde geführt wird. In der 2004 gestarteten Mission Gravity Probe B wurde diese

Bild 2:
Stephen Hawking genießt einen freien Fall an Bord einer Boeing 727.

Noted physicist Stephen Hawking (center) enjoys zero gravity during a flight aboard a modified Boeing 727 aircraft owned by Zero Gravity Corp. (Zero G). Hawking, who suffers from amyotrophic lateral sclerosis (also known as Lou Gehrig's disease) is being rotated in air by (right) Peter Diamandis, founder of the Zero G Corp., and (left) Byron Lichtenberg, former shuttle payload specialist and now president of Zero G. Kneeling below Hawking is Nicola O'Brien, a nurse practitioner who is Hawking's aide. Source: Jim Campbell/Aero-News Network/ Wikimedia Commons



sogenannte geodätische Präzession, die für die Bahn des Satelliten etwa 6,6 Bogensekunden pro Jahr beträgt, erfolgreich gemessen. Die Mission Gravity Probe B hat auch einen weiteren Effekt – wenn auch mit einiger Ungenauigkeit – bestätigt, bei dem es sich um eine der faszinierendsten Vorhersagen der ART überhaupt handelt: den Thirring-Lense Effekt. Dieser besagt, dass lokale Inertialsysteme in der Nachbarschaft eines rotierenden Körpers, hier der Erde, präzedieren, die Rotation der Erde die Bewegung der Kreiselachse also direkt beeinflusst. Im Falle des Kreisels in Gravity Probe B handelt es sich um den winzigen Effekt von etwa 41 Millibogensekunden pro Jahr.

Einstein wählte den Namen Allgemeine Relativitätstheorie, weil er eine wichtige Eigenschaft seiner Feldgleichungen zum Ausdruck bringen wollte: ihre Invarianz unter beliebigen (allgemeinen) Koordinatentransformationen. Das hat immer wieder für Verwirrung gesorgt, da Koordinaten in der allgemeinen Formulierung keine eigenständige Bedeutung mehr zukommt. Einstein selbst wurde von dieser Tatsache bei der Suche nach der Theorie zeitweise auf Irrwege geführt. Eine invariante Bedeutung hat die Eigenzeit, die für einen Beobachter auf seinem Weg durch die Raum-Zeit verstreicht, und die auf seiner Uhr angezeigt wird. Diese Eigenzeit kann durch den Gebrauch beliebiger Raum- und Zeitkoordinaten berechnet werden, ist aber unter Koordinatentransformationen invariant. Es ist die Eigenzeit, die von realen Uhren wie der Caesiumuhr angezeigt wird.

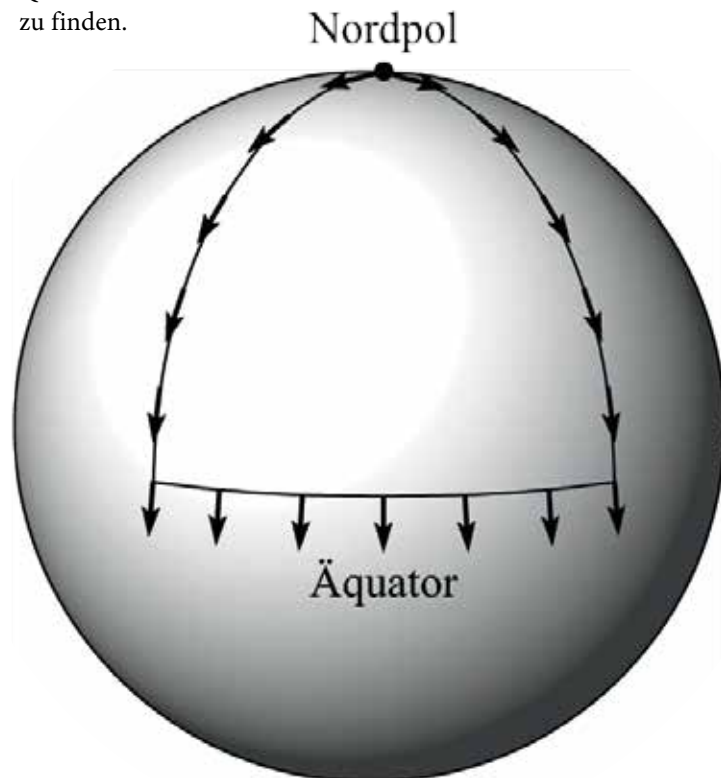
Auf diese Eigenzeit nimmt die Definition der Sekunde Bezug. Es ist deshalb verständlich, dass die Relativitätstheorie im Definitionstext nicht erwähnt wird.

Für Beobachter im gleichen Bezugssystem ist der Uhrengang unabhängig vom Bewegungszustand und vom Gravitationspotential; das folgt aus dem Äquivalenzprinzip. Nur bei Beobachtern, die eine Relativbewegung aufweisen oder sich in Gravitationspotentialen unterschiedlicher Stärke befinden, liegt ein Unterschied im relativen Uhrengang vor, der sich mit den Gesetzen von SRT und ART ausrechnen lässt (siehe unten).

Natürlich kann es in speziellen Situationen, die durch eine Symmetrie ausgezeichnet sind, ausgezeichnete Koordinaten geben. So kann man etwa bei einer sphärisch-symmetrischen Massenverteilung eine radiale Koordinate r so wählen, dass die Oberfläche zweidimensionaler Kugeln durch den euklidischen Ausdruck $4\pi r^2$ gegeben ist, obwohl r selbst nicht mehr einer geometrischen Länge entspricht.

Die prinzipielle Beliebigkeit der Koordinaten bringt zum Ausdruck, dass die Raum-Zeit in der ART keine absolute Bedeutung mehr hat. Raum und Zeit sind dynamisch und ändern sich im

Wechselspiel mit Energie- und Impulsflüssen der Materie. Da die ART Wellenlösungen kennt, die jüngst direkt nachgewiesenen Gravitationswellen, sind Raum und Zeit auch im Vakuumfall dynamisch. Es gibt keinen festen raum-zeitlichen Hintergrund mehr, eine Eigenschaft, welche die ART von allen früheren Theorien unterscheidet. Es ist auch diese Eigenschaft, die es so schwierig macht, eine Formulierung der Gravitation im Rahmen der Quantentheorie zu finden.



Die von Uhren

angezeigte Zeit – die Eigenzeit – hängt natürlich von der Stärke des Gravitationsfeldes ab. So gehen Uhren in der Nähe von Massen, wo die Krümmung stärker ist, langsamer als in deren Ferne. Eine Uhr auf dem Kölner Dom tickt also schneller als eine Uhr auf der Domplatte, wenngleich der relative Unterschied nur etwa 10^{-14} beträgt. Größer ist der Effekt für die Uhren in den Navigationssystemen GPS und Galileo. Da sich die Satelliten bewegen, kommt auch der speziell-relativistische Effekt der Zeitdilatation zum Tragen, der eine Verlangsamung bewegter Uhren bewirkt. Da gemäß der ART Uhren im All aber schneller gehen als am Erdboden, hängt es von der Höhe der Satellitenbahnen ab, welcher der beiden Effekte dominiert. Bei der Höhe der GPS-Satelliten von über 20 000 km dominiert der ART-Effekt: die Uhren gehen in bezug auf terrestrische Uhren vor, und zwar etwa um den relativen Faktor 10^{-10} . Solche Unterschiede sind mit modernen Atomuhren, deren relative Genauigkeit besser als 10^{-14} ist, leicht messbar.

Bild 3:
Verschiebt man einen Pfeil parallel vom Nordpol zum Äquator, ein Stück entlang des Äquators und dann wieder zum Nordpol, so wird er nicht mehr in die ursprüngliche Richtung zeigen; das ist ein Effekt der Krümmung der Kugeloberfläche

Die Nichtberücksichtigung der relativistischen Effekte in diesen Positionierungssystemen würde zu krassen Fehlweisungen führen. Da das Licht in einer Nanosekunde ($1 \text{ ns} = 10^{-9}$ Sekunden) etwa 30 cm zurücklegt, müssen bei einer Ortsbestimmung, die auf etwa 10 m genau sein soll, Lichtlaufzeiten auf mindestens 30 ns genau bestimmt werden können. Die relative Gangabweichung einer Satellitenuhr darf deshalb im Verlauf eines Tages (86.400 s) nicht größer als etwa 10^{-12} betragen. Da die relativistischen Effekte um etwa den Faktor 10^{-10} beitragen, würde sich deren Nichtbeachtung im Laufe des Tages auf über 8000 ns aufschaukeln, was zu einer Fehlweisung von fast 3 km führen würde; dies entspräche in 14 Tagen bereits der Entfernung Köln-Düsseldorf! Sowohl SRT als auch ART sind also Teil des Alltags geworden.

Die ART entfaltet ihre volle Kraft freilich nicht auf irdischen, sondern auf stellaren und kosmischen Skalen. Sie sagt sie etwa die Existenz von Schwarzen Löchern voraus, deren gravitative Anziehung so stark ist, dass nichts, nicht einmal Licht, aus dem Inneren ihres Ereignishorizontes entweichen kann. Auch die Wirkung auf Uhren am Rande von Schwarzen Löchern ist enorm. Während ein frei fallender Beobachter keinen Effekt bemerkt (was direkt aus dem Äquivalenzprinzip folgt), wird eine statische Uhr um so langsamer gehen (verglichen mit einer entfernten Uhr), je näher sie dem Ereignishorizont kommt. Im Grenzfall der Lage am Horizont wird sie sogar stillstehen!

Schwarze Löcher hat man bisher indirekt durch ihren Einfluss auf die Umgebung, zum Beispiel auf Akkretionsscheiben oder Sterne, nachgewiesen. Das prominenteste Beispiel ist vielleicht das supermassive Schwarze Loch Sgr A* im Zentrum unserer Milchstraße, dem man eine Masse von mehr als vier Millionen Sonnenmassen zuschreibt. Der deutlichste Nachweis von Schwarzen Löchern gelang in Verbindung mit dem ersten direkten Nachweis von Gravitationswellen. Das im September 2015 an den Ligo-Interferometern in den USA festgestellte Ereignis (das Gleiche gilt für zwei spätere Ereignisse) hat seinen Ursprung bei zwei Schwarzen Löchern in kosmischer Entfernung, die umeinander kreisten und dann verschmolzen, wobei sie die hier beobachteten Gravitationswellen aussandten. Die maximal ausgesandte Leistung entsprach mehr als 10^{49} Watt, was etwa der elektromagnetischen Leuchtkraft aller sichtbaren Sterne im beobachtbaren Universum entspricht.

Die größten Skalen, auf die Einsteins Theorie mit Erfolg angewandt wird, ist das Universum als Ganzes. Das kosmologische Standardmodell eines Universums, das aus einer heißen Frühphase expandiert und sich heute beschleunigt ausdehnt, steht mit allen bisher gemachten Beobachtungen in Einklang. Die Kosmologie zeigt freilich auch die Grenzen von Einsteins Theorie auf.

4. Zeit nach Einstein

Im Gegensatz zur Newton'schen Gravitation benennt die ART ihre eigenen Grenzen. Wie insbesondere die britischen Physiker Stephen Hawking und Roger Penrose Ende der sechziger Jahre zeigen konnten, führen sehr allgemeine Annahmen (typischerweise Energiebedingungen und die Unmöglichkeit von Reisen in die Vergangenheit) auf sogenannte Singularitäten [3]. Das sind unvermeidbare Grenzen der Raum-Zeit, an denen beispielsweise die Raumkrümmung unendlich groß werden kann. Einsteins Theorie sagt solche Singularitäten etwa für das Innere von Schwarzen Löchern und für den Anfang der Welt voraus (möglicherweise auch für deren Ende). Der Anfang der Welt, der „Zeitpunkt null“ oder „Urknall“, gehört also nicht mehr zum Anwendungsbereich der ART.

Was tun? Das Auftreten von Unendlichkeiten signalisiert in der Physik immer den Zusammenbruch einer Theorie, nie das tatsächliche Auftreten von Unendlichkeiten in der Natur. Man erwartet deshalb die Existenz einer neuen Gravitationstheorie, die Einsteins Theorie als Grenzfall umfasst. Es gibt viele Gründe, die zu der Annahme verleiten, dass es sich um eine Quantentheorie handeln müsse, eine Quantengravitation [4]. Zum einen werden alle anderen fundamentalen Wechselwirkungen durch Quanten(feld)theorien beschrieben, und eine elegante vereinheitlichte Theorie aller Wechselwirkungen einschließlich der Gravitation sollte deshalb ebenfalls eine Quantentheorie sein. Zum anderen gibt es im Verhältnis von gewöhnlicher Quantentheorie und ART ein eigenartiges Spannungsverhältnis, das man als das „Problem der Zeit“ bezeichnet. Während in Einsteins Theorie Raum und Zeit dynamisch sind, es also keinen starren Hintergrund gibt, werden die Quantentheorien der nichtgravitativen Wechselwirkungen immer auf einen solchen Hintergrund bezogen. Der Zeitparameter t in der Quantenmechanik ist gerade Newtons absolute Zeit; die Raum-Zeit in den Quantenfeldtheorien ist die feste Raum-Zeit der SRT, wie sie Minkowski eingeführt hat. Nun kann die Zeit auf fundamentaler Ebene aber nicht gleichzeitig absolut und dynamisch sein.

Solange sich die Anwendungsbereiche von Quantentheorie und ART nicht überschneiden, kommen sich die unterschiedlichen Zeitbegriffe nicht ins Gehege. Das ist aber in der Kosmologie und im Inneren der Schwarzen Löcher nicht mehr der Fall. Schwarze Löcher strahlen mit einer bestimmten Temperatur, der Hawking-Temperatur, die dem Planck'schen Wirkungsquantum proportional ist. Man benötigt deshalb eine konsistente Theorie der Quantengravitation. Doch wie sieht eine solche Theorie aus?

Es werden derzeit einige Zugänge diskutiert, ohne dass bisher eine Einigung über den korrekten Zugang zu erkennen ist. Zu diesen Zugängen gehören Stringtheorie sowie kanonische und kovariante Quantenversionen von Einsteins ART. Ein Aspekt scheint sich allerdings anzudeuten: die Zeitlosigkeit der Quantengravitation. Analog zum Verschwinden der Teilchenbahnen in der Quantenmechanik verschwindet nach der Quantisierung die Raum-Zeit und es verbleibt nur der Raum. Die Zeit ergibt sich erst wieder in einem geeigneten Grenzfall und stimmt dort mit dem Zeitbegriff der ART überein. Die Konsequenzen dieser Zeitlosigkeit insbesondere für die Kosmologie sind erst im Ansatz ausgelotet worden.

In seinem 1953 verfassten Vorwort zu Max Jammers Buch über den Begriff des Raumes schrieb Einstein die folgenden Worte:

„Es hat schweren Ringens bedurft, um zu dem für die theoretische Entwicklung unentbehrlichen Begriff des selbständigen und absoluten Raumes zu gelangen. Und es hat nicht geringerer Anstrengung bedurft, um diesen Begriff nachträglich wieder zu überwinden – ein Prozeß, der wahrscheinlich noch keineswegs beendet ist.“ [7]

Natürlich gilt das Gesagte genauso für die Zeit. Ohne die Begriffe absoluter Raum und absolute Zeit hätte Newton die Gleichungen der klassischen Mechanik und sein Kraftgesetz der universellen Gravitation nicht formulieren können. Und ohne die Überwindung dieser Begriffe wäre Einstein nicht zur Relativitätstheorie gelangt. Die Vereinigung dieser Theorie mit der Quantentheorie führt zu einem Zeitbegriff, wie ihn sich weder Newton noch Einstein vorstellen konnten.

Literatur:

- [1] H. Blumenberg, *„Die Genesis der kopernikanischen Welt“*, (Frankfurt am Main 1975)
- [2] D. Giulini, *„Spezielle Relativitätstheorie“*, (Frankfurt am Main 2004)
- [3] S. Hawking und R. Penrose, *„Raum und Zeit“*, (Reinbek 2000)
- [4] C. Kiefer, *„Der Quantenkosmos“*, (Frankfurt am Main 2008)
- [5] J. Schwinger, *„Einsteins Erbe“*, (Heidelberg 1987)
- [6] J. Stachel (Hg.), *„Einsteins Annus mirabilis“*, (Reinbek 2001); enthält Einsteins Originalarbeiten von 1905
- [7] M. Jammer, *„Das Problem des Raumes“* (Darmstadt 1960)

Auf dem Weg zu einer Neudefinition der Sekunde: Was kommt nach Caesium?

Ekkehard Peik*

Große Fortschritte in der Entwicklung der sogenannten optischen Uhren lösten bereits vor einigen Jahren die Diskussion über eine Neudefinition der Sekunde aus. Unterschiedliche Prototypen solcher Uhren sind in mehreren Laboratorien entwickelt worden, und die besten Experimente haben dabei eine Stabilität und Genauigkeit der Frequenz erreicht, die primäre Caesiumuhren um etwa zwei Größenordnungen übertrifft. Schon mehrmals ließen sich erhebliche Verbesserungen in der Technik der Zeitmessung an der Einführung einer neuen Generation von Uhren festmachen, die jeweils bei deutlich höherer Taktfrequenz arbeitete als ihre Vorgänger. Angefangen von Sonnenuhren, Sand- oder Wasseruhren (mHz), über mechanische Uhren mit Pendeln oder Federn (Hz), zu Quarzuhren (kHz), den ersten Atomuhren (GHz), bis hin zu den optischen Atomuhren, die mit einer Frequenz von sichtbarem oder ultraviolett Licht (10^{15} Hz) arbeiten.

Optische Uhren:

Eine neue Generation von Atomuhren

Mit der Anhebung der Taktfrequenz um fünf Größenordnungen von der Caesiumuhr zur optischen Uhr sind deutliche Gewinne bei Frequenzstabilität und Genauigkeit zu erreichen. So lässt sich die Frequenzstabilität bei kurzen Messzeiten durch die wesentlich höhere Zahl von Schwingungszyklen pro Zeiteinheit verbessern. In der Genauigkeit erzielt man Gewinne, da manche der störenden äußeren Effekte auf das Atom eine charakteristische Größenordnung an Verschiebung ΔE der Energieniveaus hervorrufen, die als relative Frequenzverschiebung $\Delta f/f = \Delta E/hf$ bei der höheren Frequenz der optischen Uhr einen kleineren Effekt als für eine Uhr im Mikrowellenbereich bewirkt.

Als Oszillator einer optischen Uhr dient stets ein Laser. Die Entwicklung von Cs-Atomuhr, Maser und Laser erfolgte etwa gleichzeitig in den 1950er-Jahren und die daran arbeitenden Forscher hatten ähnliche wissenschaftliche Erfahrungen und Interessen im Bereich der Spektroskopie von

Atomen und Molekülen. So war zum Beispiel Harold Lyons, der 1948 am National Bureau of Standards (heute NIST) in den USA eine Atomuhr mit Ammoniak-Molekülen entwickelt hatte (s. den Aufsatz von M. Lombardi in diesem Heft), später Leiter der Arbeitsgruppe der Hughes Research Laboratories, in der Theodore Maiman 1960 erstmals die Erzeugung von Laserlicht gelang. Frühe Veröffentlichungen über Maser und Laser erwähnen die Realisierung von verbesserten Frequenznormalen bereits als eine der ersten absehbaren Anwendungen der neuartigen Strahlungsquellen [1].

Wichtige Beiträge zu den heutigen optischen Uhren wurden in den 1970er- und 1980er-Jahren mit der Entwicklung von Methoden zur Speicherung und Laserkühlung von Atomen und Ionen in Fallen geleistet. Die Lokalisierung und Reduzierung der Bewegungsenergie von Atomen ist notwendig, um Frequenzverschiebungen durch den Doppler-Effekt zu vermeiden, der wegen der kürzeren Wellenlänge im optischen Spektralbereich schwieriger zu kontrollieren ist als bei Mikrowellen. Außerdem ermöglichen in Fallen gespeicherte Atome beliebig lange Wechselwirkungszeiten mit dem Laserlicht, und damit eine verbesserte Frequenzauflösung. Eine Falle lässt sich besonders gut für geladene Atome (Ionen) mit elektrischen Feldern realisieren. Durch die elektrostatische Kraft auf die Ladung lässt sich das Ion fixieren, ohne dabei seine innere Struktur, also die für die Uhr entscheidende Resonanzfrequenz, wesentlich zu stören. Die ersten Vorschläge für eine optische Uhr wurden von Hans Dehmelt in den 1970er-Jahren gemacht und beruhen auf einem einzelnen, in einer Ionenfalle gespeicherten und gekühlten Ion [2]. Auf neutrale Atome lassen sich Kräfte nur dadurch ausüben, dass man an der Ladungsverteilung im Atom angreift und damit zwangsläufig die elektronische Struktur beeinflusst. Erst 2001 wurde es durch einen Vorschlag von Hidetoshi Katori [3] klar, dass man auch mit neutralen Atomen eine sehr genaue optische Uhr bauen kann, indem man eine Falle konstruiert, die beide Energieniveaus

* Dr. Ekkehard Peik, Fachbereich „Zeit und Frequenz“, E-Mail: ekkehard.peik@ptb.de

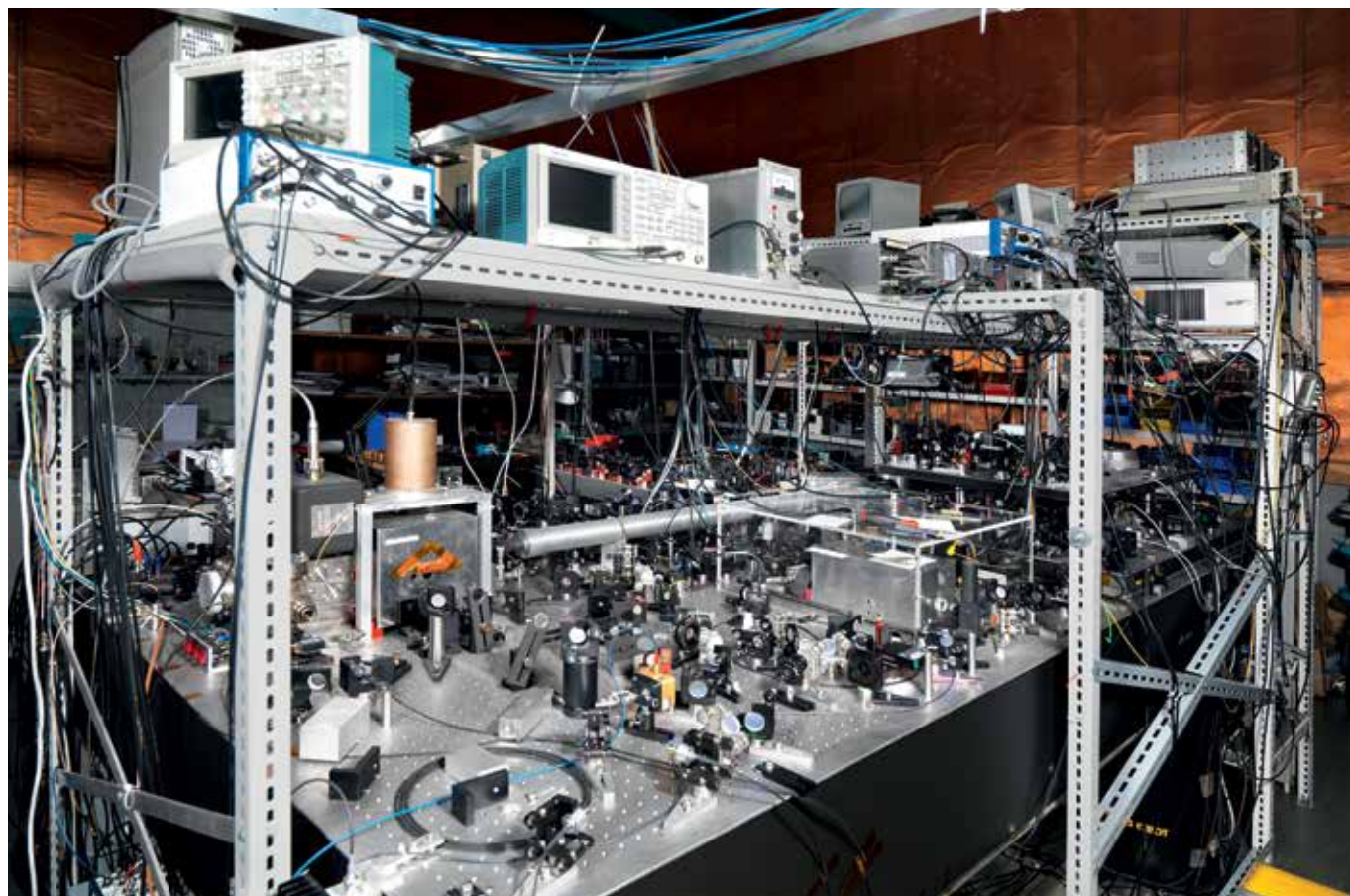
des Übergangs, der die Referenzfrequenz der Uhr bestimmt, genau gleich verschiebt. Die Falle wird aus dem Interferenzmuster von Laserstrahlen gebildet, deren „magische“ Wellenlänge so gewählt ist, dass die genannte Bedingung erfüllt ist. Man nennt diese Anordnung ein optisches Gitter, da sich viele, die Atome einschließende Potentialminima periodisch im Raum bilden. Hier ergibt sich ein wichtiger Vorteil einer Uhr mit neutralen Atomen: Es lassen sich gleichzeitig einige Tausend Atome kühlen und speichern und durch die Struktur der Fallen werden Stöße zwischen den Atomen, die die beobachtete Resonanzfrequenz verschieben könnten, weitgehend vermieden. Bis heute werden beide Konzepte als „Einzel-Ionenuhr“ bzw. „Gitteruhr“ weiterentwickelt [4] (s. Bild 1).

Die Methoden der Speicherung und Laserkühlung sind auf eine Vielzahl von Atomen und Ionen in unterschiedlichen Ladungszuständen anwendbar. Bei der Auswahl des Atoms für eine optische Uhr spielen daher die Eigenschaften des Referenzübergangs eine wichtige Rolle. In den Arbeiten von Dehmelt wird auf das Ion Tl^+ hingewiesen, und diese Auswahl beruht auf vorteilhaften Eigenschaften seiner Elektronenhülle: Im Grundzustand und auch im ersten angeregten Zustand koppeln die zwei Valenzelektronen zu einem Paar mit verschwindendem Drehimpuls. Dadurch sind die von äußeren elektrischen und magnetischen Feldern hervorgerufenen Energieverschiebun-

gen deutlich kleiner als in Zuständen mit einem einzelnen Valenzelektron. Der Übergang zwischen beiden Niveaus verletzt die Auswahlregeln für elektrische Dipolstrahlung, und besitzt wegen der langen Lebensdauer des angeregten Zustands die für eine Atomuhr benötigte geringe natürliche Linienbreite. Wegen anderer praktischer Nachteile wurden bisher keine Experimente zu einer Tl^+ -Uhr gemacht, aber die meisten der heute für optische Uhren untersuchten Elemente besitzen dieselbe Struktur mit zwei Valenzelektronen: die neutralen Atome Sr, Hg, Yb und die Ionen In^+ und Al^+ .

Basierend auf diesen grundlegenden Konzepten wurden in den zurückliegenden Jahren spektakuläre Fortschritte in der Entwicklung optischer Uhren erreicht, angefangen von spektroskopischer Auflösung von Linienbreiten im Bereich von einem Hertz bei optischen Übergängen, über Frequenzmessungen mit Caesiumuhren im Unsicherheitsbereich 10^{-15} bis hin zu systematischen Unsicherheiten im Bereich 10^{-18} . Damit sind optische Uhren jetzt auf kurzen Zeitskalen etwa um einen Faktor 100 stabiler und in der Reproduzierbarkeit ihrer Frequenz auch um einen Faktor 100 genauer als primäre Caesiumuhren. Das Ziel einer Genauigkeit von 10^{-18} war von Dehmelt erstmals 1981 konkret benannt worden, damals um acht Größenordnungen jenseits der besten realisierten Atomuhren. Bis heute, Mitte 2017, haben fünf Gruppen in den USA, Japan und Deutschland mit

Bild 1:
Ansicht des Aufbaus der optischen Uhr mit einem einzelnen gespeicherten Yb^+ -Ion an der PTB.



unterschiedlichen Experimenten, zwei mit Ionen Al^+ - und Yb^+ - und drei mit Sr- und Yb-Gitteruhren, diesen Bereich erreicht.

Die Sekunde im Neuen SI-System

Deutlich dringlicher als bei der SI-Einheit „Sekunde“ ist das Bestreben nach neuen Definitionen bei den SI-Einheiten Ampere, Kilogramm und Kelvin. Bei den elektrischen Einheiten sind auf Quanteneffekten basierende Realisierungen von Widerstand und Spannung besser reproduzierbar als die auf der bisherigen SI-Definition basierenden Normale, und bei der Masse-Einheit besteht Sorge über die langfristige Stabilität des internationalen Prototypen (das „Urkilogramm“ in Paris) (s. die Beiträge in PTB-Mitteilungen Heft 2/2016). In den vorgesehenen Neudefinitionen sollen die sieben bisherigen Basiseinheiten des SI über die festgelegten Zahlenwerte von sieben Naturkonstanten definiert werden: so zum Beispiel das Ampere über die Elementarladung e und das Kilogramm über die Planck'sche Konstante h . Diese umfassende Reform des SI-Systems, die 2018 beschlossen werden soll, folgt dem bereits 1983 für die Definition des Meters (durch Festlegung des Wertes der Lichtgeschwindigkeit c) angewandten Modell. Ein Ziel dieser Reform ist es auch, die Definitionen der Einheiten unabhängig von spezifischen, praktischen Realisierungen zu halten, um so in der Zukunft flexibel technische Fortschritte in den Realisierungen für die Anwendung im SI-System umsetzen zu können. Für die Sekunde ergibt sich hierbei zunächst keine Änderung: Die definierende Konstante bleibt die Frequenz des Caesium-Hyperfeinstrukturübergangs mit 9 192 631 770 Hz, und damit bleibt die Sekunde eng mit ihrer Realisierung, der Caesiumuhr, verknüpft.

Die Sekunde erhält eine zentrale Rolle im neuen SI, da sie bei sechs der sieben definierenden Konstanten in deren zusammengesetzten Einheiten enthalten ist. Die einzige Ausnahme bildet die Avogadro-Konstante, als natürliche Zahl, die die Einheit der Stoffmenge Mol definiert. Es ist daher ein wichtiger Aspekt, dass sich die SI-Sekunde mit der mit Abstand geringsten Unsicherheit aller SI-Basiseinheiten realisieren lässt, wodurch ihr Beitrag zur Unsicherheit der Realisierung anderer Einheiten praktisch vernachlässigbar ist.

Optionen für eine neue Definition der SI-Sekunde

Rückblickend betrachtet kann man die SI-Definition der Sekunde von 1967 als ein Erfolgsmodell bezeichnen: Eine geeignete atomare Übergangsfrequenz wird unter idealen Bedingungen realisiert und eine festgeschriebene Zahl von Schwingungs-

perioden als Zeiteinheit festgelegt. Es erscheint daher naheliegend, dieses Modell für die Zukunft mit einem geeignet gewählten optischen Übergang weiterzuführen. Nun eröffnet die Diskussion um eine Neudefinition aber auch die Möglichkeit, über Alternativen nachzudenken, und auch dies wurde in den letzten Jahren getan. Zwei Vorschläge sollen hier kurz diskutiert werden [5].

Eine an den theoretischen Grundlagen der Physik orientierte Option wäre es, die Definition der Sekunde auf einem berechenbaren Frequenznormal aufzubauen, auf einer Frequenz, die sich in einfacher Weise durch fundamentale Naturkonstanten angeben lässt.¹ Geeignet erscheint dafür am ehesten die Rydbergfrequenz $cR_\infty = (m_e e^4) / (8 \epsilon_0^2 h^3) \approx 3,2898 \cdot 10^{15}$ Hz, berechnet aus der Masse m_e und Ladung e des Elektrons, der elektrischen Feldkonstanten ϵ_0 und der Planck-Konstante h . Diese einfache Formel für die Rydbergfrequenz ergibt sich aus dem Bohr'schen Atommodell, und sie beschreibt die (nicht-relativistische) Skalenkonstante für alle Übergangsfrequenzen in der Elektronenhülle von Atomen und Molekülen. Experimentell präzise bestimmt – mit einer relativen Unsicherheit von einigen 10^{-12} – wurde die Rydbergfrequenz aus Übergangsfrequenzen von atomarem Wasserstoff. Wären also Wasserstoff, oder Wasserstoff-ähnliche Ionen wie He^+ , Li^{2+} etc, als berechenbare Uhren die besten Elemente für die Realisierung einer optischen Sekundendefinition? Grundsätzliche und auch praktische Probleme stehen dem entgegen: So liefert die Bohr'sche Theorie nur eine näherungsweise Beschreibung der Übergangsfrequenzen. Beiträge der Relativitätstheorie und Quantenelektrodynamik müssen berücksichtigt werden und können nur näherungsweise berechnet werden. Nicht berechenbare Einflüsse der inneren Struktur des Protons tragen in der Größenordnung 10^{-12} zu den Übergangsfrequenzen von Wasserstoff bei, weit oberhalb der für eine optische Uhr angestrebten Genauigkeitsskala. Auch sind die Techniken der Präzisionsspektroskopie mit gespeicherten und gekühlten Atomen für Wasserstoff bisher nicht vergleichbar weit entwickelt, wie für die schwereren Atome, die heute in optischen Uhren verwendet werden.

Gibt es vielleicht andere, aus den bekannten Fundamentalkonstanten ableitbare, universelle Frequenzen? Eine einfache Betrachtung der Einheiten macht klar, dass Kombinationen der im SI-System verwendeten Konstanten h (Einheit $\text{m}^2\text{kg/s}$), c (Einheit m/s) und e^2/ϵ_0 (Einheit $\text{m}^3\text{kg/s}^2$) es nicht erlauben, eine Größe der Einheit Sekunde oder Hertz zu beschreiben. Hierfür ist zusätzlich eine Konstante mit der Einheit kg, also z. B. die Masse eines Elementarteilchens, erforderlich. Verwendet man die Elektronenmasse, so findet man die Rydbergfrequenz oder die

¹ Dies ist weder bei der Caesiumuhr noch bei den heute untersuchten optischen Uhren der Fall: Die atomare Struktur dieser Elemente ist so komplex, dass numerische Berechnungen von Übergangsfrequenzen nur die ersten 4–6 Stellen der experimentellen Ergebnisse reproduzieren können.

Comptonfrequenz $\nu_C = m_e c^2 / h \approx 1,236 \cdot 10^{20}$ Hz des Elektrons. Rydberg- und Comptonfrequenz sind über $cR_\infty = \alpha^2 \nu_C / 2$ (mit $\alpha \approx 1/137$ der Feinstrukturkonstanten) miteinander verknüpft. Die Comptonfrequenz des Elektrons liegt bereits weit oberhalb des heute für Präzisionsmessungen zugänglichen Bereichs. Die Verwendung der Masse eines schwereren Teilchens würde die Frequenz zu noch höheren Werten verschieben. Für die Idee eines primären Frequenznormals mit berechenbarer Frequenz ist somit heute keine praktikable Realisierung absehbar.

Da derzeit mehrere, auf unterschiedlichen atomaren Übergängen basierende Frequenznormale ähnlich geringe Unsicherheiten erreichen, wurde vorgeschlagen, anstelle *eines* atomaren Übergangs mit festgelegter Frequenz, *mehrere*, gleichberechtigt für eine neue Definition der Sekunde zu verwenden. Diese Vorgehensweise würde die möglicherweise kontroversen Diskussionen um die Auswahl eines spezifischen Übergangs abmildern, und später den nationalen Metrologieinstituten und anderen Anwendern eine Auswahl von Realisierungen der SI-Sekunde eröffnen. Es gibt aber absehbare Schwachpunkte dieser Lösung: Durch die Festlegung von mehreren Zahlenwerten wäre die Definition überbestimmt. Mit Weiterentwicklungen in den Realisierungen der ausgewählten Frequenzen würde dies eine Motivation schaffen, durch verbesserte Messungen von Frequenzverhältnissen die festgelegten Konstanten immer wieder neu anzupassen. Dies würde es erforderlich machen, alle präzisen Frequenzmessungen mit ihrer Historie bezüglich der verwendeten Frequenzreferenz und des Umrechnungsfaktors zu dokumentieren. Um die Konsistenz des Systems langfristig zu wahren, könnte allen festgelegten Frequenzwerten eine geringe Unsicherheit zugeordnet werden. Diese könnte so klein sein, dass sie auf die Messungen anderer, von der Sekunde abhängender Einheiten keinen Einfluss hätte. Das würde aber die Paradoxie ergeben, dass gerade die SI-Einheit, die sich mit der geringsten Unsicherheit realisieren lässt, bereits eine Unsicherheit in ihrer Definition tragen würde.

Die 1967 gefundene Struktur der SI-Definition der Zeiteinheit erscheint somit auch nach heutigem Stand des Wissens als so zweckmäßig, dass sie sich für die Fortschreibung mit einem neuen Referenzübergang anbietet. Zur Wahrung der Kontinuität der Einheit wäre lediglich der neue Zahlenwert der Schwingungsperioden pro Zeiteinheit experimentell aus Messungen des Frequenzverhältnisses der optischen Uhr zur Cesiumuhr zu bestimmen.

Auf dem Weg zur Neudefinition

In Vorbereitung eines solchen Verfahrens werden seit 2001 von einer Arbeitsgruppe des beratenden Komitees für Zeit und Frequenz CCTF die Ergebnisse von Frequenzmessungen geeigneter Übergänge gesammelt, evaluiert und daraus empfohlene Frequenzen für sogenannte sekundäre Realisierungen der Sekunde abgeleitet [6]. Zurzeit werden neun verschiedene Übergänge empfohlen: die in einer Rubidiumfontäne realisierbare Hyperfeinstrukturfrequenz von ^{87}Rb , die in optischen Gittern untersuchten Systeme ^{87}Sr , ^{171}Yb und ^{199}Hg und fünf Übergänge der Ionen $^{27}\text{Al}^+$, $^{88}\text{Sr}^+$, $^{171}\text{Yb}^+$ (mit zwei Übergängen) und $^{199}\text{Hg}^+$. In den kommenden Jahren werden mehr und mehr optische Frequenzverhältnisse zwischen diesen und anderen Referenzübergängen gemessen werden. Als dimensionslose Zahlen lassen sie sich unabhängig von der Cs-basierten Unsicherheit der SI-Sekunde bestimmen und können zwischen den Laboratorien, die gleiche Paare von optischen Uhren betreiben, verglichen und für Konsistenztests im Unsicherheitsbereich von 10^{-18} verwendet werden.

Parallel dazu wird die Technologie optischer Uhren weiterentwickelt, um sie einsatzbereit für die Anwendungsbereiche heutiger Atomuhren zu machen, etwa im Dauerbetrieb zur Steuerung einer Zeitskala und als kommerziell erhältliches, schlüsselfertiges Gerät für Laboratorien, die keine eigene Entwicklungsarbeit in diesem Gebiet betreiben. Hier wird das kürzlich gestartete Quantentechnologie-Flaggschiffprojekt der EU wichtige Impulse geben.

Die Vielfalt der aussichtsreichen Kandidaten und die schnell fortschreitende Entwicklung neuer Techniken und Methoden erlauben es heute noch nicht, einen klaren Favoriten für die Nachfolge von Cesium auszumachen. Vermutlich werden, wie es heute im Mikrowellenbereich mit Cesiumuhren, Rubidiumuhren und Wasserstoffmasern der Fall ist, je nach Anwendung unterschiedliche optische Uhren eingesetzt werden. Da die weitaus meisten Anwendungen von Atomuhren und atomaren Frequenznormalen heute im Mikrowellenbereich liegen und die derzeitigen Anforderungen dort gut erfüllen, werden Cesiumuhren auf absehbare Zeit nicht an Bedeutung verlieren.

Es erscheint also verfrüht, jetzt bereits die Auswahl des Referenzübergangs in Richtung auf eine baldige Neudefinition der Sekunde einzuschränken. Es ist dagegen als gutes Zeichen für die wissenschaftliche Dynamik des Feldes zu sehen, dass sich die Rangfolge der jeweils „besten“ Uhren immer wieder ändert. Weitere, neuartige Systeme wurden vorgeschlagen und kämpfen derzeit noch mit experimentellen Startschwierigkeiten, werden aber möglicherweise innerhalb weniger Jahre in

die Spitzengruppe aufschließen. Beispiele für neue Referenzsysteme sind ein auf dem niederenergetischen Isomer in Th-229 basierender optischer Kernübergang und Übergänge in gespeicherten, hochgeladenen Ionen.

Die interessantesten Anwendungen der neuen Uhren liegen derzeit in der Wissenschaft, insbesondere bei der Suche nach neuer Physik an den Grenzen der grundlegenden Theorien, wie der Relativitätstheorie und der Quantenphysik. Durch die unterschiedliche Struktur (wie Masse und Kernladung) der verwendeten atomaren Referenzsysteme würden sich bisher nur vermutete mögliche Effekte wie zeitliche Änderungen der Kopplungskonstanten der vier fundamentalen Wechselwirkungen oder eine Kopplung an Dunkle Materie in Frequenzverschiebungen zwischen den unterschiedlichen Atomuhren äußern. Die hohe Präzision in der Messung von Frequenzen könnte damit ein Fenster zu neuen Einblicken in die Grundlagen der Physik öffnen.

Referenzen:

- [1] C. H. Townes, „*How the laser happened*“, Oxford Univ. Press, 1999
- [2] H. Dehmelt, „*Coherent spectroscopy on single atomic system at rest in free space*“, *J. Phys.* (Paris) **42**, C8–299 (1981)
- [3] H. Katori, „*Spectroscopy of Strontium atoms in the Lamb-Dicke confinement*“, in *Proc. of the 6th Symposium on Frequency Standards and Metrology*, World Scientific, Singapore, 2002
- [4] A. D. Ludlow, M. M. Boyd, Jun Ye, E. Peik, P. O. Schmidt, „*Optical Atomic Clocks*“, *Rev. Mod. Phys.* **87**, 637 (2015)
- [5] P. Gill, „*When should we change the definition of the second*“, *Phil. Trans. R. Soc. A* **369**, 4109 (2011)
- [6] F. Riehle, „*Towards a redefinition of the second based on optical clocks*“, *C. R. Physique* **16**, 506 (2015)

Empfangssystem für glasfasergeführtes ultrapräzises Frequenzsignal

Eine PTB-Erfindung ermöglicht die Übertragung des Signals eines ultrastabilen Single-Frequency-Lasers über große Entfernungen in normalen Telekommunikationsglasfasern. Die Erfindung löst das Problem des Anschlusses einer großen Anzahl von Kunden an eine einzige Faserstrecke. Die Erfindung stellt zugleich einen bedeutenden Schritt zur Übertragung des Zeitsignals einer optischen Uhr dar (Atomuhr aus der Steckdose).

Technische Beschreibung

Eine optische Frequenz wird über eine lange Glasfaserleitung übertragen und kann jetzt – trotz der zu erwartenden Störungen in Phase, Mittenfrequenz und Polarisation – an jedem Ort der Leitung abgetastet und auf die Ursprungsfrequenz von des PTB-Normals rückgeführt werden. Dies gelingt, indem beide Signale, sowohl das vorwärtslaufende als auch das rückwärtslaufende, zu einem Schwebungssignal vereinigt werden. Ein nachfolgender einfacher, analoger Algorithmus erzeugt eine Korrekturfrequenz $\Delta\nu$. Ein akustooptischer Modulator (AOM) überlagert nun das gestörte Signal mit dieser Korrekturfrequenz und regeneriert das gewünschte PTB-Frequenzsignal von in hoher Präzision. Im Empfänger werden einfache Standardkomponenten der Telekommunikationstechnik eingesetzt. Nachfolgende Empfangsstationen werden durch die Auskopplung am Faserkoppler einer einzelnen Station nicht gestört. In Punkt-zu-Punkt-Experimenten ist die Übertragungstechnik über Wegstrecken von mehr als 100 km Länge nachgewiesen worden.

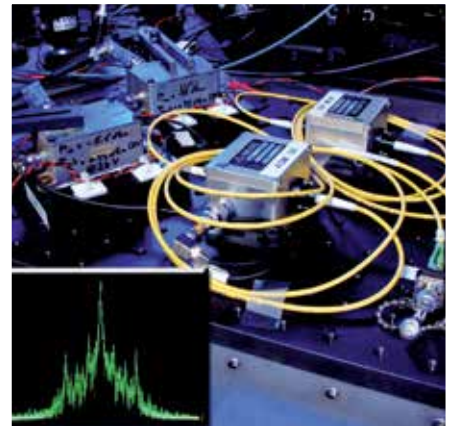
Wirtschaftliche Bedeutung

Empfänger dieser Art werden im Bereich der Lasertechnik bei der Kalibrierung von höchstauflösenden Spektrometern und der Feinabstimmung von lokalen Referenzlasern auf Empfängerseite benötigt. Kombiniert mit einem lokalen Frequenzkammgenerator können weitere präzise Frequenzen, auch im Mikrowellenbereich, erzeugt werden. Im Bereich der chemischen Analytik mit Höchstpräzisionslasern kann das System mittelbar zur Verifikation der Nachweisempfindlichkeit und Qualitätssicherung dienen.

Eine lokale Verteilung kann mit nur geringen Qualitätseinbußen verschiedene Arbeitsplätze bedienen. Durch viele Empfangsstationen reduziert sich der Mietpreis der benötigten normalen, überall vorhandenen 1,5- μ m-Glasfaserleitung.

Entwicklungsstand

Unter DE 10 2008 062 139 B4 wurde ein Patent erteilt.



Die Erfindung ermöglicht den Anschluss vieler Empfänger an eine einzige Glasfaserstrecke. Jeder Empfänger ist in der Lage, trotz Übertragungsstörungen das ursprüngliche hochpräzise Frequenzsignal zu regenerieren.

Vorteile:

- Bereitstellung des Frequenzsignals mit einer Bandbreite um 1 Hz in einem normalen Glasfaserkabel über große Entfernungen
- Auf Normal rückführbar
- Mehrfache Abtastung an einem beliebigen Ort
- Verteilung des Frequenzsignals in lokalen Netzen über einige hundert Meter mit geringen Qualitätseinbußen

Ansprechpartner:

Dr. Bernhard Smandek
Beauftragter für Technologietransfer
Telefon: +49 531 592-8303
Telefax: +49 531 592-69-8303
E-Mail: bernhard.smandek@ptb.de

Dr. Gesine Grosche
Arbeitsgruppe
Frequenzübertragung mit Glasfasern



Optisches Phasenrausch-Messgerät im Laboraufbau

Vorteile:

- Phasenrauschen im sub-Hz-Bereich wird messbar
- Einfaches Verfahren mit Standardkomponenten
- Erprobter optischer Aufbau

Optisches Phasenrausch-Messgerät

Viele Lasersysteme erreichen mittlerweile Linienbreiten im Bereich weniger kHz, sodass deren Phasenrauschen nur schwer zu charakterisieren ist. Ohne weiteren Zusatzlaser generiert die PTB-Erfindung mittels einer hochstabilen externen Kavität das Bezugssignal für die Messung selbst.

Technische Beschreibung

Um das verbleibende Phasenrauschen hochstabiler Laser messen zu können, wird das Licht üblicherweise im Selbstheterodynverfahren über eine lange Glasfaser mit sich selbst verglichen. Mittlerweile sind zum Beispiel Faserlaser in der Telekommunikation jedoch so gut, dass die Selbstheterodynmethode nicht mehr empfindlich genug ist. Die Verwendung eines zweiten noch rauschärmeren Lasers ist aus Kostengründen häufig unpraktikabel.

Die PTB-Methode ist in der Abbildung dargestellt. Sie besteht darin, ein Referenzsignal aus der Quelle selbst mittels einer externen Kavität zu generieren und dies als hochgenaue Referenzquelle zu nutzen. Das Rauschspektrum wird dann üblicherweise über die Charakterisierung der Schwebungsfrequenzen bestimmt. Hiermit sind Messungen des Frequenzrauschens für Fourierfrequenzen von 10 MHz bis zu einigen MHz möglich. Im Bereich von 1–10 Hz erreicht das Verfahren eine Messempfindlichkeit unterhalb von 10^{-2} Hz²/Hz.

Wirtschaftliche Bedeutung

Der neuartige Spiegel kann in Fabry-Perot-Resonatoren eingesetzt werden, die zur Stabilisierung von Lasersystemen dienen. Solche Lasersysteme erreichen eine relative Frequenzstabilität von bis zu 10^{-16} . Hochstabile Laser mit geringer Linienbreite finden in der Präzisionspektroskopie, der Fasersensorik für Baustatik oder Geologie und in der Telekommunikation ihre Anwendung.

Ansprechpartner:

Dr. Bernhard Smandek
Beauftragter für Technologietransfer
Telefon: +49 531 592-8303
Telefax: +49 531 592-69-8303
E-Mail: bernhard.smandek@ptb.de

Dr. Thomas Legero
Fachbereich
Quantenoptik und Längeneinheit

Entwicklungsstand

Das Verfahren wird routinemäßig in der PTB zur Lasercharakterisierung eingesetzt. Unter DE 10 2012 024 692 B3 wurde ein Patent erteilt.

Spiegelbauteil für ultrastabile Resonatoren

Ein neuartiges Spiegel-Konzept für Fabry-Perot-Resonatoren führt durch eine entscheidende Verringerung des thermischen Rauschens zu einer höheren Frequenzstabilität des Resonators bei geringerem Aufwand für die Temperaturstabilisierung. Die Kombination aus Design und Materialauswahl führt zu einer relativen Frequenzstabilität im Bereich von 10^{-16} . Dabei lässt sich das Spiegelbauteil einfach und kostengünstig fertigen.

Technische Beschreibung

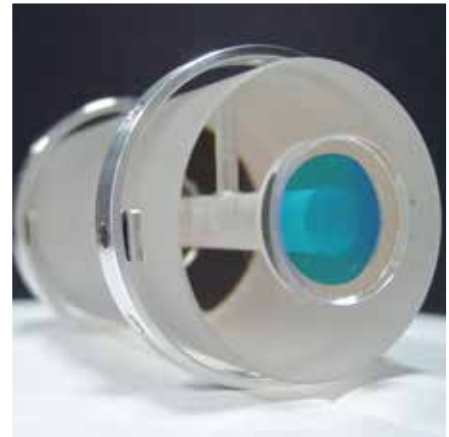
Das neue Spiegelement der PTB kombiniert eine extrem hohe Längenstabilität bei ca. 20 °C mit geringem thermischen Rauschen. Bisher bestanden sowohl die Spiegel als auch der Abstandhalter von ultrastabilen Fabry-Perot-Resonatoren aus Ultra-Low-Expansion-Glas (ULE), das bei ca. 20 °C einen Nulldurchgang in seiner Längenausdehnung aufweist. Bei der PTB-Lösung wird an der Rückseite eines Quarzglasspiegels ein Ring aus ULE-Glas angebracht. Dieser wirkt der Deformation des Spiegels entgegen, und der Nulldurchgang der Längenausdehnung des Resonators befindet sich wieder bei 20 °C. Damit sind Fabry-Perot-Resonatoren möglich, die bei Zimmertemperatur betrieben werden können und dennoch ein deutlich kleineres thermisches Rauschen aufweisen.

Wirtschaftliche Bedeutung

Der neuartige Spiegel kann in Fabry-Perot-Resonatoren eingesetzt werden, die zur Stabilisierung von Lasersystemen dienen. Solche Lasersysteme erreichen eine relative Frequenzstabilität von bis zu 10^{-16} . Frequenzstabile Laser werden zur Realisierung optischer Uhren, in der Nachrichtentechnik zum Übermitteln von ultrastabilen Frequenzen und im Bereich der Ultrapräzisions-Spektroskopie benutzt.

Entwicklungsstand

Für das Spiegelbauteil wurde unter DE 10 2008 049 367 B3 ein Patent erteilt.



Ultrastabiler optischer Resonator mit einem Längenausdehnungskoeffizienten von Null bei ca. 20 °C

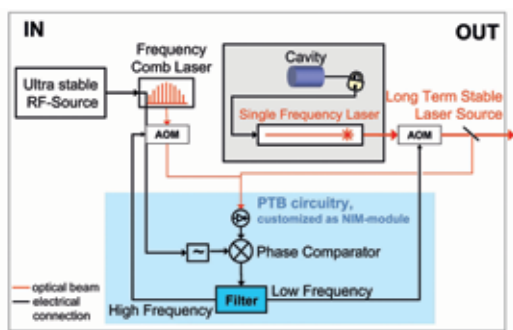
Vorteile:

- Einsatz von Spiegelmaterialien mit geringerem thermischen Rauschen
- Kompensation der durch Temperaturschwankungen verursachten Spiegeldeformation
- Relative Längenstabilität von 10^{-16} möglich
- Kostengünstige Temperaturstabilisierung auf 20 °C

Ansprechpartner:

Dr. Bernhard Smandek
Beauftragter für Technologietransfer
Telefon: +49 531 592-8303
Telefax: +49 531 592-69-8303
E-Mail: bernhard.smandek@ptb.de

Dr. Thomas Legero
Fachbereich
Quantenoptik und Längeneinheit



Schematische Darstellung des Wirkprinzips

Langzeitstabiler Single-Frequency Laser

Schmalbandige Laser mit extrem kleinen Bandbreiten finden zunehmend Anwendung in der Spektroskopie, der Metrologie, der Astrophysik und der Ultrapräzisionsmesstechnik. Hier werden inzwischen Linienbreiten ≤ 1 Hz routinemäßig erreicht. Die Langzeitstabilität solcher Systeme war bisher gar nicht oder nur mit unverhältnismäßigem Aufwand zu realisieren. Durch die patentierte Kopplung an eine Radioquelle hoher Stabilität bietet die PTB hierfür jetzt eine technische Lösung.

Technische Beschreibung

Die Grundidee besteht darin, die vorteilhaften Stabilitätseigenschaften zweier Frequenzquellen zu kombinieren, beispielsweise die Kurzzeitstabilität (= schmale Linienbreite) eines hochstabilen Lasers mit langzeitstabilen Quellen, die typischerweise im RF-Bereich emittieren. Beispiele hierfür sind ein Wasserstoff-Maser, eine Chip-Caesium-Atomuhr oder auch einfach ein GPS-disziplinierter Oszillator.

Durch geeignete Kombination der sich ergänzenden Eigenschaften verschiedener Frequenzquellen, kann so eine Frequenzquelle realisiert werden, die im optischen Bereich eine extrem schmale Linienbreite mit einer hohen Langzeitstabilität von z. B. besser als $5 \cdot 10^{-15}$ über lange Zeiten kombiniert.

Wirtschaftliche Bedeutung

Eine schmalbandige, optische Frequenzquelle mit hoher Langzeitstabilität eröffnet die Möglichkeit von Präzisionsmessungen über lange Zeiträume, beispielsweise in der wissenschaftlichen Spektroskopie, der Astrophysik oder auch in Analyselaboren in verschiedenen Wirtschaftsbereichen. Weiterhin ist sie gut geeignet zum Vermessen und Charakterisieren schmalbandiger Lichtquellen, beispielsweise in der photonischen Industrie.

Vorteile:

- Langzeitstabile optische Frequenz
- Sub-Hz-Linienbreite
- Verwendung bewährter kommerzieller Komponenten
- PTB-Elektronik-Einschübe im NIM-Format

Ansprechpartner:

Dr. Bernhard Smandek
Beauftragter für Technologietransfer
Telefon: +49 531 592-8303
Telefax: +49 531 592-69-8303
E-Mail: bernhard.smandek@ptb.de

Dr. Sebastian Raupach
Arbeitsgruppe
Frequenzübertragung mit Glasfasern

Entwicklungsstand

Die spezielle Regelungseinheit wurde im NIM-Einschubformat für 19-Zoll-Racks entwickelt und in der PTB testweise eingesetzt. Eine abschließende Entwicklung ist durch den Lizenznehmer alleine oder in Kooperation mit der PTB möglich. Unter DE 10 2012 008 456 B4 wurde ein Patent erteilt.

Impressum

Die PTB-Mitteilungen sind metrologisches Fachjournal der Physikalisch-Technischen Bundesanstalt, Braunschweig und Berlin. Als Fachjournal veröffentlichen die PTB-Mitteilungen wissenschaftliche Fachaufsätze zu metrologischen Themen aus den Arbeitsgebieten der PTB. Die PTB-Mitteilungen stehen in einer langen Tradition, die bis zu den Anfängen der Physikalisch-Technischen Reichsanstalt (gegründet 1887) zurückreicht.

Verlag

Fachverlag NW in der
Carl Schünemann Verlag GmbH
Zweite Schlachtpforte 7
28195 Bremen
Internet: www.schuenemann.de
E-Mail: info@schuenemann-verlag.de

Herausgeber

Physikalisch-Technische Bundesanstalt (PTB)
ISNI: 0000 0001 2186 1887
Postanschrift:
Postfach 33 45,
38023 Braunschweig
Lieferanschrift:
Bundesallee 100,
38116 Braunschweig

Redaktion/Layout

Presse- und Öffentlichkeitsarbeit, PTB
Sabine Siems
Dr. Dr. Jens Simon (verantwortlich)
Dr. Andreas Bauch
(wissenschaftlicher Redakteur)
Telefon: (05 31) 592-82 02
Telefax: (05 31) 592-30 08
E-Mail: sabine.siems@ptb.de

Leser- und Abonnement-Service

Karin Drewes
Telefon (0421) 369 03-56
Telefax (0421) 369 03-63
E-Mail: drewes@schuenemann-verlag.de

Anzeigenservice

Karin Drewes
Telefon (0421) 369 03-56
Telefax (0421) 369 03-63
E-Mail: drewes@schuenemann-verlag.de

Erscheinungsweise und Bezugspreise

Die PTB-Mitteilungen erscheinen viermal jährlich.
Das Jahresabonnement kostet 39,00 Euro, das Einzelheft 12,00 Euro, jeweils zzgl. Versandkosten.
Bezug über den Buchhandel oder den Verlag.
Abbestellungen müssen spätestens drei Monate vor Ende eines Kalenderjahres schriftlich beim Verlag erfolgen.

Alle Rechte vorbehalten. Kein Teil dieser Zeitschrift darf ohne schriftliche Genehmigung des Verlages vervielfältigt oder verbreitet werden. Unter dieses Verbot fällt insbesondere die gewerbliche Vervielfältigung per Kopie, die Aufnahme in elektronische Datenbanken und die Vervielfältigung auf CD-ROM und in allen anderen elektronischen Datenträgern.

Printed in Germany ISSN 0030-834X

Die fachlichen Aufsätze aus dieser Ausgabe der PTB-Mitteilungen sind auch online verfügbar unter:
doi: 10.7795/310.20170399



Bundesministerium
für Wirtschaft
und Energie

Die Physikalisch-Technische Bundesanstalt, das nationale Metrologieinstitut, ist eine wissenschaftlich-technische Bundesoberbehörde im Geschäftsbereich des Bundesministeriums für Wirtschaft und Energie.

Helmholtz Preis 2018

SIE HABEN SICH DIE **METROLOGIE** IN DEN KOPF GESETZT?

Zur Förderung von Wissenschaft und Forschung verleihen der Helmholtz-Fonds e. V. und der Stifterverband für die Deutsche Wissenschaft e. V. gemeinsam den Helmholtz-Preis für außergewöhnliche Leistungen auf dem Gebiet der Präzisionsmessung.

Er wird in diesem Jahr in zwei Kategorien ausgeschrieben:

1. Präzisionsmessung in der Grundlagenforschung der Physik, Chemie und Medizin
2. Präzisionsmessung in der angewandten Messtechnik der Physik, Chemie und Medizin

Der Preis besteht in jeder Kategorie aus einer Urkunde und einem Preisgeld von 20.000 €.

Weitere Infos und Bewerbungseinreichung (bis zum 18.12.2017) über:
www.helmholtz-fonds.de/helmholtz-preis



HELMHOLTZ
FONDS e.V.

www.helmholtz-fonds.de