

Aus den Elfenbeintürmen der Wissenschaft 5
XLAB Science Festival

Aus den Elfenbeintürmen der Wissenschaft 5

XLAB Science Festival

*Herausgegeben von
Eva-Maria Neher*



WALLSTEIN VERLAG

Die XLAB Science Festivals 2010 und 2011
wurden vom Universitätsbund Göttingen e.V.
und vom DFG-Forschungszentrum
Molekularphysiologie des Gehirns (CMPB) gefördert.

Mitarbeiter und Weggefährten Manfred R. Schroeders
stifteten zu seinem Andenken die *Manfred-Schroeder-Lecture*
beim Science Festival 2011.

Der Druck des Bandes
wurde durch die Klaus Tschira Stiftung ermöglicht.

Klaus Tschira Stiftung
Gemeinnützige GmbH



Inhalt

EVA-MARIA NEHER	
Wissenschaft ist international	7
JULIA FISCHER	
Zum Ursprung der menschlichen Sprache	15
A. JULIA STÄHLER	
Unsterbliche Elektronen	39
MARKUS HARTL, JAN-W. KELLMANN UND IAN T. BALDWIN	
Ökologische Gentechnik	47
CHRISTOPH MARKSCHIES	
Wilhelm und Alexander	85
JENS FRAHM, MARTIN UECKER, DIRK VOIT UND SHUO ZHANG	
Magnetresonanz-Tomografie in Echtzeit	98
ULF DIEDERICHSEN	
Lineare molekulare Architekturen mit DNA- und Peptid-Biooligomer-Gerüsten	III
HANNS RUDER UND HANS-PETER NOLLERT	
Was Einstein gern gesehen hätte	124
ULRICH CHRISTENSEN UND NORBERT KRUPP	
Die Geschwister der Erde	151
NINA SCHALLER	
Auf Zehenspitzen zum Weltrekord	166

INHALT

AARON CIECHANOVER (Nobelpreis 2004)	
Intracellular Protein Degradation: From a Vague Idea thru the Lysosome and the Ubiquitin-Proteasome System and onto Human Diseases and Drug Targeting	175
ECKART ALTENMÜLLER, OLIVER GREWE, FREDERIK NAGEL UND REINHARD KOPIEZ	
Musik als Sprache der Gefühle	207
BIRGER KOLLMEIER	
Partys, MP ₃ -Player, Hörgeräte: Leistungen physikalischer Hörmodelle	216
Die Autoren.	239
Farbabbildungen	143

Eva-Maria Neher

Wissenschaft ist international

International Science Camps am XLAB

Die Wissenschaften selber werden wohl weder als international noch als national bezeichnet werden können. Vielmehr werden die wissenschaftlichen Fragestellungen international von vielen Wissenschaftlern bearbeitet. Die Wissenschaften sind auch nicht national, sondern können national begrenzte Themen umfassen oder, wie die Geschichte immer wieder lehrt, politisch oder von nationalistischem Gedankengut oder Zielen geprägt sein. Es sind also die Fragestellungen und das Interesse an ihrer Lösung entscheidend dafür wo, von wem und in welcher Kooperation zu einem Thema geforscht wird.

In Europa war Griechenland die Geburtsstätte der Wissenschaften, die von vielen naturwissenschaftlichen Kenntnissen aus Ägypten und Babylonien geprägt war. Eine frühe Entwicklung der Wissenschaften gab es auch im Chinesischen Reich. Die ersten wissenschaftlichen Interessen galten der Astronomie, der Mathematik und in China auch der Medizin. In der westlichen Welt gab es keine Trennung zwischen den Naturwissenschaften und der Philosophie, wohingegen in China praktische und technisch orientierte Fragestellungen gar nicht als eine Aufgabe der Gelehrten angesehen wurden. Wissenschaftliche Erkenntnis lebt davon, dass sie weitergegeben und somit weiterentwickelt werden kann. Die Überlieferung von Wissen erfolgte in den Anfängen mündlich, später durch die Schrift. Die älteste bekannte Schrift ist die sumerische Keilschrift, sie entstand etwa um 3500 v. Chr. Unabhängig entwickelt sich etwa 1000 v. Chr. in China und in der Kultur der Maya in Mittelamerika eine Schrift. Geschrieben wurde auf Materialien, die entweder wenig dauerhaft oder oftmals nicht transportabel waren. Für die Verbreitung von Wissenschaft ist aber die Verfügbarkeit dauerhafter und leicht transportierbarer Informationsträger eine unabdingbare Voraussetzung. Die Erfindung des Papiers ist ein Verdienst der Chinesen. Die erste Beschreibung der chinesischen Methode zur Herstellung von Papier datiert aus dem Jahr 105 n. Chr. In Europa wurden die Methoden der Papierherstellung verfeinert und mechanisiert, so

dass zusammen mit der von Johannes Gutenberg (1398-1468) entwickelten Technik des Buchdrucks ein neues Zeitalter für die Verbreitung von Wissen begann.

Die ältesten abendländischen Universitäten waren Bologna, gegründet Ende des 11. Jh., gefolgt von Paris, Oxford, Cambridge, Neapel, Avignon, Pisa, Prag, Krakau, Wien und anderen im 13. und 14. Jh. Die Gründung der ersten deutschen Universitäten erfolgte an der Wende vom 14. zum 15. Jh. (Heidelberg 1386, Würzburg 1402, Leipzig 1409). Im Zuge des spanischen Kolonialismus in Mittel- und Südamerika kam es zu ersten Universitätsgründungen, beispielsweise in Peru (1551), Mexiko (1553) und vielen anderen Ländern. In Nordamerika fanden die ersten Universitätsgründungen erst im 18. Jahrhundert statt, so in Yale (1701), gefolgt von Princeton (1746), New York (1754), Philadelphia (1755) und Washington (1789).

Die Expeditionen der Seefahrer, z.B. die Entdeckung des Seeweges nach Indien durch Vasco da Gamas (1460-1524), die Ankunft von Christoph Kolumbus (1451-1506) in Amerika, die Weltumseglung von Ferdinand Magellan (1480-1561) oder die Entdeckungsreisen von James Cook (1728-1779) im Pazifischen Ozean, brachten mehr und mehr Kenntnisse über die Geographie der Erde. Sie schufen die Grundvoraussetzungen für Reisen in andere Kontinente und machten die Erforschung fremder Kulturen möglich. Herausragende Beispiele für Naturforscher, die auf umfangreichen und beschwerlichen Reisen ihren Forschungsinteressen auf der ganzen Welt folgten, sind Alexander von Humboldt (1769-1859)¹ oder Charles Darwin (1809-1882). Was sie trieb, war die Neugier; ihr Wissensdurst ließ sie die Strapazen ertragen. Unterwegs tauschten sie sich mit Gelehrten der Gastländer aus. Obwohl solche Expeditionen nur wenigen privilegierten Persönlichkeiten, die mit genügend Geld und Empfehlungsschreiben der Herrscherhäuser ausgestattet waren, vorbehalten waren, können sie als Anfang einer globalen Entwicklung der Wissenschaft betrachtet werden.

Als älteste deutsche Akademie der Wissenschaften wurde die Leopoldina 1652 von vier Ärzten in der Freien Reichsstadt Schweinfurt

1 Vgl. Christoph Markschie's Artikel »Wilhelm und Alexander. Das Verhältnis Alexanders zu seinem älteren Bruder Wilhelm von Humboldt« in diesem Band.

gegründet. In den Gründungsstatuten wird »die weitere Aufklärung auf dem Gebiet der Heilkunde und der daraus hervorgehende Nutzen für die Mitmenschen« zur Richtschnur für das Handeln der Akademie erklärt.² Die Leopoldina verstand sich ausdrücklich als übernationale Gelehrtenvereinigung. Im 19. Jahrhundert erfolgte die Gründungen weiterer Fachgesellschaften, deren grundlegendes Ziel es war, Foren für den Gedankenaustausch innerhalb der wissenschaftlichen Gemeinschaft zu eröffnen. So ist beispielsweise die Deutsche Physikalische Gesellschaft e.V. (DPG), deren Tradition bis in das Jahr 1845 zurückreicht, die älteste nationale und mit über 60 000 Mitgliedern auch größte physikalische Fachgesellschaft der Welt. Etwa zwanzig Jahre später erfolgte 1867 in Berlin die Gründung der Gesellschaft Deutscher Chemiker (GDCH). Die GDCH zählt heute 30 000 Mitglieder und ist in internationale Netzwerke eingebunden. Die älteste mir bekannte internationale wissenschaftliche Fachgesellschaft ist die *International Union of the Physiological Societies* (IUPS). Die IUPS ist der weltweite Zusammenschluss vieler nationaler physiologischer Akademien. Sie fand sich 1889 in Basel zusammen und veranstaltet seitdem alle drei Jahre internationale Kongresse. Erst 1953 fand das erste Treffen in einem außereuropäischen Land in Montreal, Kanada statt. Von da an war der Weg frei für Kongresse auf dem gesamten Erdball: Buenos Aires (1959), Tokio (1965), Washington (1968), Neu-Delhi (1974), Sydney (1983), Vancouver (1986), St. Petersburg (1997), Christchurch (2001), San Diego (2005), Kyoto (2009) und zukünftig Rio de Janeiro (2017), dazwischen liegen viele hier nicht angeführte Tagungsorte in Europa. Das weltumspannende Netzwerk war natürlich gekoppelt an die Entwicklung der zivilen Luftfahrt. Die Fluggastzahlen vom Flughafen Frankfurt am Main betrugen 1960 noch zwei Millionen, 2010 bereits 53 Millionen jährlich. Nur so ist es möglich geworden, mehrere tausend Wissenschaftler aus 70 und mehr Ländern zusammenzuführen. Ich hatte das Vergnügen, seit 1997 einige der Tagungsorte besuchen zu dürfen. Es war beeindruckend, die Wissenschaftler aus vielen Ländern treffen zu können. Viele kennen einander sehr gut und tauschen regelmäßig ihre Ergebnisse aus. Wen wundert es, dass sie einander auch freundschaftlich zugetan sind. Doktoranden und Postdoktoranden ge-

2 Wiedergegeben nach der Homepage der Leopoldina <http://www.leopoldina.org/de/akademie/geschichte.html> am 18.7.2011.

hören dazu und wachsen in diese wissenschaftliche Gemeinschaft hinein. Die Erfahrung dieser *Community* hat mich sehr in dem Wunsch bestärkt, Ähnliches auch für deutlich jüngere, zukünftige Nachwuchswissenschaftler einzurichten. In diesem Ansinnen hat mich Manfred Eigen, Nobelpreisträger und Schirmherr des XLAB Science Festivals, stets unterstützt (Farbabbildung 1, S. 143).

Interkulturelle Austauschprogramme für Jugendliche gibt es viele, schon in der Schulzeit. Sommerprogramme an den im Ausland sehr verbreiteten *Boarding schools* oder *Colleges* finden in der unterrichtsfreien Zeit statt und sind oft auch internationalen Teilnehmern zugänglich. Auch wenn ein naturwissenschaftlicher Schwerpunkt gesetzt wird, überwiegen sportliche und kulturelle Aktivitäten, wofür an den *Colleges* alle Voraussetzungen gegeben sind. Naturwissenschaftliches Lernen bestimmt hingegen an der *Summer Science School* am *Weizmann Institute of Science* in Rehovot, Israel, am *Petnica Science Center* in Serbien oder bei den *International Science Camps* am Göttinger XLAB den Tagesablauf. Die Programme in Israel und Serbien sind entweder in der Form von Laboraufenthalten in verschiedenen wissenschaftlichen Abteilungen organisiert oder durch Diskussionsrunden und Exkursionen geprägt.

Die *International Science Camps* im XLAB bieten den jungen Teilnehmerinnen und Teilnehmern ein kursmäßig organisiertes Lernen in allen Naturwissenschaften, das von einem Kultur- und Freizeitprogramm begleitet wird. In den kleinen, auf 12 Teilnehmer begrenzten Kursen wird Teamarbeit erlernt. Sie sind geprägt durch Exaktheit beim Experimentieren in den forschungsnah ausgestatteten Laboren, die Konzentration auf ein Thema und Englisch als Unterrichts- und Diskussionssprache. So bietet das XLAB im Jahr 2011 zum 15. Mal *International Science Camps* von dreieinhalb Wochen Dauer an. Jede Woche wird ein anderer Experimentalkurs in den Fachbereichen Biologie, Chemie oder Physik absolviert. Die Kursthemen sind u.a. Anatomie, Embryologie, Pflanzenphysiologie, Grüne Gentechnik, Wald- und Gewässerökologie, Molekularbiologie, Chemie der Arzneimittel, Analytische Chemie, Laserphysik, Wellenphysik oder die Physik des Fliegens. Am Ende einer jeden Kurswoche werden die Erfahrungen und Ergebnisse im Plenum vorgestellt. An den Wochenenden gibt es Exkursionen. In der vierten Woche wird das Camp mit einer Fahrt nach Berlin abgeschlossen, die dann auch landeskundliche und gesellschaftspolitische Inhalte hat. *Science Camps* dieser Art sind weltweit

einzigartig. Wen wundert es, dass der Zuspruch groß ist. Doch es gibt begrenzende Faktoren: Zum einen akzeptieren wir nur vier Teilnehmer eines Landes, um sicherzustellen, dass sich keine muttersprachlichen Untergruppen formieren. Die gemeinsame Sprache ist die Sprache der Wissenschaft und die heute allgemein verbindliche Basis ist Englisch. Bei einer Kursgröße von 12 Teilnehmern können wir nur eine maximale Zahl von 36 Teilnehmern annehmen. Das ist eine anspruchsvolle Relation von Teilnehmern zu Dozenten, aber nur so ist ein hohes Maß an Qualität garantiert. Darauf legen wir größten Wert. Natürlich wäre es uns möglich, bei gleicher Qualität mehr parallele Kurse anzubieten, doch können wir noch nicht mehr Teilnehmer adäquat unterbringen. Der zweite begrenzende Faktor ist finanzieller Natur. Der Teilnehmerbeitrag von derzeit 1850 Euro deckt nur etwa 50 % der Kosten für den Aufenthalt. Manche der Teilnehmer werden durch nationale Stipendien gefördert, viele – besonders die deutschen – aber nicht. Den Preis zu erhöhen, hieße, eine noch höhere Hürde zu setzen. Wir aber wollen die engagierten, naturwissenschaftlich höchst motivierten jungen Menschen ansprechen und nicht nur diejenigen, die aus begüterten Elternhäusern kommen.

Die Notwendigkeit eines Gästehauses für das XLAB kann nicht oft genug unterstrichen werden. Das ganze Jahr über fehlt den Kursteilnehmern des XLAB eine preisgünstige Unterkunft auf dem Nordcampus der Universität Göttingen. Es mangelt an Gemeinschaftsräumen, ganz zu schweigen von einem Raum, in dem die jungen Menschen beispielsweise musizieren oder sich in angenehmer Atmosphäre austauschen können. Die Jugendlichen, die den Weg ins *XLAB International Science Camp* gefunden haben, sind nicht nur naturwissenschaftlich hoch interessiert, sondern vielseitig begabt und Botschafter ihrer Länder und ihrer jeweiligen Kultur.

450 junge Menschen haben von 2003 bis 2011 an den *International Science Camps* teilgenommen. Sie und einige weitere aus anderen internationalen Programmen des XLAB sind Mitglieder der *XLAB Alumni Association* geworden. Zu allen halten wir engen Kontakt und verfolgen ihre wissenschaftliche Laufbahn. Für viele war die Teilnahme am *XLAB International Science Camp* ein Sprungbrett in die Welt. Nicht wenige haben in der Folge an renommierten Universitäten in Amerika oder England studiert oder sind heute in Deutschland in einem Master- und PhD-Programm. Wir sind stolz auf unsere Alumni.

Im Zuge des so genannten Bologna-Prozesses wird die gegenseitige Anerkennung einzelner im Ausland absolvierter Studienmodule angestrebt. Auch hier konnte das XLAB erfolgreich aktiv werden. Im Rahmen zweier Kooperationen mit den Universitäten in Sevilla und Granada kommen jährlich bis zu 30 spanische Medizinstudenten im 2. Studienjahr nach Göttingen, um am XLAB drei einwöchige Kurse in den Fächern Anatomie, Molekularbiologie und Neurophysiologie zu absolvieren. Dafür werden ihnen an ihrer Heimatuniversität *Credits* angerechnet. Warum kommen diese Studierenden ans XLAB? Die Antwort ist einfach: Im Medizinstudium in Spanien ist eine experimentelle Ausbildung in den ersten Jahren nicht vorgesehen. Man verspricht sich von dieser Praxiserfahrung, mehr Medizinstudenten für die Forschung begeistern zu können. Bei uns werden Organe des Schweins präpariert, von denen viele den menschlichen Organen anatomisch sehr ähnlich sind. Wichtiger als das Studium der Anatomie der Organe ist für die Studenten, den Umgang mit dem Skalpell zu üben, feine Strukturen wie Gefäße, Nerven, Sehnen u.a. zu erkennen und zu präparieren. Einblick in die menschliche Anatomie ist an Körperspenden gebunden und bleibt als Herz der medizinischen Ausbildung selbstverständlich universitären Einrichtungen vorbehalten.

Deutschland ist derzeit ein sehr beliebtes Ausbildungsland für wissenschaftliche Nachwuchskräfte. Politisch ist das Werben um Ausländer gewollt. So ist im Jahresbericht 2010 der Max-Planck-Gesellschaft zu lesen, dass 33 % der Wissenschaftler und sogar 52 % der Nachwuchs- und Gastwissenschaftler eine ausländische Staatsangehörigkeit haben. Dies zeigt den hohen Grad der Internationalisierung unserer Forschungseinrichtungen. Als Wissenschaftssprache ist Englisch an deutschen Instituten – in zunehmendem Maß auch an den Universitäten – vorherrschend.

Bei allen Bemühen um Internationalisierung im Lande darf nicht vergessen werden, dass wir eine hohe Verantwortung für die jungen Menschen unseres Landes haben. Wir dürfen sie als wichtige Zukunftsträger nicht vernachlässigen. Auch sie brauchen die frühe Erfahrung der Internationalität von Wissenschaft, d.h. Lernen und Arbeiten im international zusammengesetzten Team. Wenn ich immer wieder davon spreche, das Konzept »XLAB« in andere Länder transportieren zu wollen, so habe ich damit die Absicht, möglichst vielen jungen Deutschen eine wissenschaftliche Lernerfahrung im Ausland zu ermöglichen. Die Gründung des *Network of Youth Excellence* (NYEX

e.V.)³ unter Beteiligung von Partnern in Israel, Serbien, Ungarn, Spanien, Italien, Korea, Rumänien und mehreren deutschen Einrichtungen ist eine Keimzelle für die Bildung eines globalen Netzwerks mit dem Ziel der internationalen Begegnung von jungen zukünftigen Nachwuchswissenschaftlern. Ein deutscher Teilnehmer in unserem ersten *Science Camp* 2003, der sich von seiner Heimatstadt noch kaum entfernt hatte, betonte, dass dies ja nicht notwendig sei, da das XLAB ihm ermögliche, viele fremde Kulturen zu Hause kennen zu lernen. Deutlich anders würden wir unseren Anspruch formulieren, ihm und den anderen Teilnehmern der *International Science Camps*, den Weg in die internationale Wissenschaft zu bahnen und Auslandsaufenthalte schon zu Beginn des Studiums zu ermöglichen. Das XLAB ist für junge Menschen aller Nationalitäten ein Tor zur Welt.

3 NYEX wurde vor etwa zehn Jahren in Ungarn von Professor Péter Csermely gegründet. Im November 2009 wurde NYEX e.V. nach dem in Deutschland gültigen Vereinsrecht als gemeinnützige Körperschaft mit Sitz in Göttingen registriert.

Julia Fischer

Zum Ursprung der menschlichen Sprache

Eine evolutionsbiologische Perspektive¹

*The roots of language remain buried,
the language gap wide and unexplained.*

Joel Wallmann

Der Ursprung der Sprache gilt als eine der großen Fragen der Menschheit, die bereits in den verschiedenen Schöpfungsmythen aufgegriffen und mit Beginn der Aufklärung auch Gegenstand der forschenden und kritischen Betrachtung wurde. So fragte im Jahre 1769 die Berliner Akademie der Wissenschaften: »Sind die Menschen, wenn sie ganz auf ihre natürlichen Fähigkeiten angewiesen sind, imstande, die Sprache zu erfinden? Und mit welchen Mitteln gelangen sie aus sich heraus zu dieser Erfindung?« Der erste Preis wurde Johann Gottfried Herder für seine »Abhandlung über den Ursprung der Sprache« (Herder 1772) zugesprochen. Herder nahm an, dass die ersten Worte den Lauten von Tieren nachempfunden gewesen seien. Anderen Überlegungen zufolge sind die ersten Wörter aus emotionalen Ausdrücken der Freude, der Angst, oder des Schmerzes entstanden. Eine weitere Theorie, die sich derzeit wieder großer Beliebtheit erfreut, rankt sich um die Annahme, dass die früheste menschliche Verständigung nicht aus Lauten, sondern aus Handzeichen und Gebärden bestand (Kuckenburg 2004). Diese drei wesentlichen Erklärungsansätze tauchten bis ins 19. Jahrhundert hinein in immer verschiedenen Varianten und Kombinationen in der Diskussion auf, bis sich die linguistische Gesellschaft von Paris 1866 genötigt sah, der ausufernden Debatte Einhalt zu gebieten und jedwede weitere Spekulation über das Thema zu verbieten.² Damit be-

- 1 Julia Fischer: Zum Ursprung der menschlichen Sprache. Eine evolutionsbiologische Perspektive. In: Neuronen im Gespräch – Sprache und Gehirn, hrsg. v. Helmut Fink und Rainer Rosenzweig, Paderborn (mentis Verlag), S. 141 ff. Abdruck mit freundlicher Genehmigung des mentis Verlags.
- 2 Keine aktuelle Veröffentlichung zum Thema des Sprachursprungs kommt ohne diesen Hinweis aus.

gründete sie die eigentliche linguistische Forschung, da man sich nun mit der Struktur von Sprache beschäftigen konnte.

Seit einigen Jahrzehnten gilt die Frage nach dem Ursprung der Sprache wieder als der wissenschaftlichen Auseinandersetzung zugänglich. Zum einen lieferten verschiedene Disziplinen wie die Paläoanthropologie, die Neurobiologie und die Verhaltensforschung reichlich Material, über das nun debattiert werden konnte. Außerdem bietet die evolutionsbiologische Forschung einen Ansatz, um sich der Frage nach dem Ursprung der Sprache auf produktive Weise zu nähern. Mit Hilfe des vergleichenden Ansatzes steht eine etablierte Methode zur Verfügung, sich dem Sprachursprung aus evolutionsbiologischer Perspektive zu nähern. Dabei werden kommunikative Fähigkeiten von in der heutigen Zeit lebenden Tieren untersucht und Vergleiche mit unseren eigenen kommunikativen Fähigkeiten angestellt, aus denen sich dann Rückschlüsse ziehen lassen, welche Eigenschaften wann und in welchem Zusammenhang entstanden sein könnten. Ziel ist, herauszufinden, welche Eigenschaften homolog sind und auf gemeinsame Vorfahren zurückgeführt werden können, welche spezifisch für nur ein Taxon³ oder eine Art sind, und welche letztlich Konvergenzen darstellen, also aufgrund ähnlicher selektiver Drücke entstanden sind. Auch Herder bediente sich eines ähnlichen Ansatzes, denn er schrieb: »Wo finge der Weg der Untersuchung sichrer an als bei Erfahrungen über den Unterschied der Tiere und Menschen?« (Herder 1772). Dieser Frage sei im Folgenden nachgegangen, wobei ich mich zunächst auf die vokal-auditorische Kommunikation bei unseren nächsten Verwandten, den nichtmenschlichen Primaten (nachfolgend Primaten), beziehen werde. Die Kernthese, die ich in diesem Aufsatz vertrete ist, dass es eine fundamentale Dichotomie in der Flexibilität in der Lautgebung im Vergleich zum Verständnis von Lauten gibt: Während die Struktur der Lautmuster weitgehend angeboren ist, spielt Lernen bei der Entwicklung der adäquaten Reaktionen und der Zuweisung von Bedeutung zu Lautmustern eine entscheidende Rolle. Wie offen und flexibel das System des Lautlernens sein kann, werde ich am Beispiel des Wortlernens bei Haushunden darlegen. Schließlich soll diskutiert werden, welche

3 Taxon (griech. *tattein* = ordnen). Als Taxon (das; Plural *Taxa*) bezeichnet man in der Biologie eine als systematische Einheit erkannte Gruppe von Lebewesen.

Faktoren bei der Entwicklung der Fähigkeit, Lautmuster zu imitieren, eine Rolle gespielt haben können.

Die Untersuchung der linguistischen Fähigkeiten von Primaten begann in den 50er Jahren des 20. Jahrhunderts mit einer Reihe von Studien an in Gefangenschaft gehaltenen, zum Teil in Familien aufgezogenen Schimpansen. Zum Beispiel versuchten Cathy und Keith Heyes in den 1950er Jahren, der Schimpansin Vicki gesprochene Sprache beizubringen. Allerdings beschränkte sich Vickis Vokabular auch nach intensivem Training lediglich auf die Begriffe *mama*, *up* und *cup*, wobei sie die beiden letztgenannten Worte nur flüstern konnte und ihre Hand zur Hilfe nehmen musste, um den Explosivlaut /p/ zu produzieren. Die meisten späteren Studien basierten daher entweder auf dem Einsatz von Gebärdensprache oder von abstrakten Symbolen, die auf Computermonitoren oder Tafeln angeboten wurden. Die natürlichen kommunikativen Fähigkeiten von Affen fanden erst ein Jahrzehnt später Beachtung, als sich Louis Leakey und Sherwood Washburn mit ihrer Ansicht durchsetzten, dass sich die evolutionären Ursprünge des menschlichen Verhaltens nur durch das Studium der nächsten lebenden Verwandten verstehen ließen. Bezogen auf die Frage nach der Evolution der menschlichen Sprache rückten damit die Lautäußerungen von Primaten ins Zentrum der Aufmerksamkeit (Übersicht in Wallmann 1992).

Bei der Untersuchung der Lautgebung hat es sich bewährt, zwischen drei Aspekten des kommunikativen Verhaltens zu unterscheiden, und zwar der Lautproduktion, dem Einsatz von Lauten und dem Verständnis von Lautmustern. Fragen nach der Entwicklung der Lautproduktion zielen primär darauf ab, altersabhängige Veränderungen in der Struktur der Lautmuster zu identifizieren; allgemeiner umfassen sie aber auch den Aspekt der Modifizierbarkeit von Lautmustern als Folge sozialer Erfahrung. Untersuchungen der Entwicklung des Einsatzes von Lautmustern gehen der Frage nach, wie Primaten mit zunehmendem Alter und Erfahrung verschiedene Lautmuster in unterschiedlichen Kontexten differentiell äußern. Veränderungen im Einsatz von Lautmustern sind im Allgemeinen ein Resultat einer verfeinerten Kategorisierung und Bewertung von sozialen Situationen und externen Stimuli. Solche Untersuchungen geben daher vornehmlich Einblick in die Entwicklung der kognitiven Leistungen der Tiere. Weiterhin ist von Interesse, wie sich die Reaktionen auf die Lautmuster von Artgenossen, aber auch von anderen Tieren entwickeln. Dies beinhaltet

zugleich die Frage, wie Tiere lernen, Lautmuster mit Bedeutung zu belegen, und die Frage, ab welchem Alter adäquate Reaktionen zu beobachten sind.

Ontogenese der Lautproduktion bei nichtmenschlichen Primaten

Der überwiegende Teil der Studien zur Ontogenese der Lautproduktion bei nichtmenschlichen Primaten legt den Schluss nahe, dass die Struktur der Lautmuster weitgehend angeboren ist. Modifikationen treten nur in einem sehr begrenzten Umfang auf, und Erfahrung mit arteigenen Lautmustern ist in der Regel keine Voraussetzung für die Entwicklung eines vollständigen und normalen Lautrepertoires. Dennoch verändert sich die Struktur von Lautmustern mit dem Alter, und sie unterscheidet sich auch hinsichtlich des Geschlechts. Dies sei am Beispiel der Kontaktrufe von Bärenpavianen dargestellt (Abb. 1). Kontaktrufe eignen sich für eine Analyse altersabhängiger Unterschiede, da sie zu den wenigen Ruftypen gehören, die von Tieren aller Altersklassen geäußert werden. Wie beim Menschen wird die Stimme im Verlauf des Alters tiefer, was sich im Absenken der Grundfrequenz, aber auch der Verschiebung der Energie im Frequenzspektrum äußert (Ey et al. 2007b). Dies lässt sich vor allem auf das Größenwachstum zurückführen, mit dem eine Verlängerung der Stimmbänder einhergeht. Akustische Variablen, die die Struktur der Lautmuster charakterisieren, verändern sich hingegen nicht. Allerdings sind nicht alle Änderungen lediglich mit körperlichem Wachstum zu erklären. So finden sich bei jungen Tieren kaum Unterschiede zwischen den Geschlechtern. Mit dem Einsetzen der Pubertät hängen die Männchen eine zweite, etwas leisere Silbe an den *wa*-Laut an – dementsprechend werden die Lautäußerungen der adulten Männchen auch als *wa-boos* bezeichnet. Diese zweite Silbe setzt in etwa dann verstärkt ein, wenn auch der Testosteronspiegel ansteigt und die männlichen Tiere weitere sekundäre sexuelle Merkmale entwickeln wie zum Beispiel längere Eckzähne und prächtiges Schulterhaar. Ranghohe Männchen sind in der Lage, sehr viel lautere und längere zweite Silben zu produzieren als rangniedere Tiere. Außerdem verliert die zweite Silbe binnen weniger Monate an Ausdruck, wenn ein Tier im Rang gefallen ist. Letzten Endes, wenn die Männchen alt und gebrechlich sind, hören sich ihre Rufe wieder an

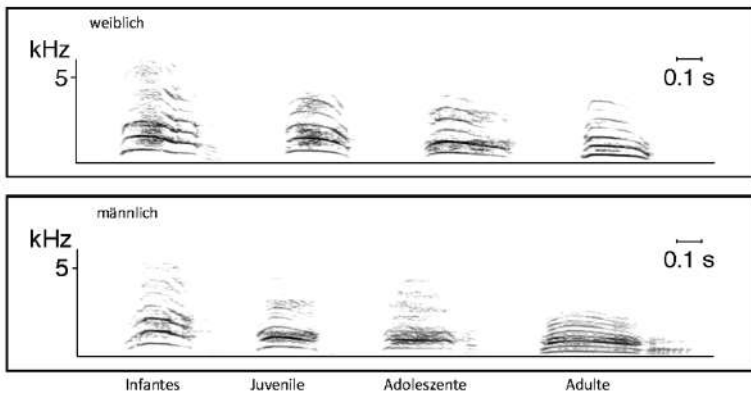


Abbildung 1: Spektrographische Darstellung der Kontaktrufe von Bärenpavianen. Auf der Y-Achse ist die Frequenz, auf der X-Achse die Zeit dargestellt. Unterschiedliche Schwärzungsgrade geben die Energieverteilung wieder: Dunklere Bereiche entsprechen Frequenzen mit hoher Energie, hellere Bereiche solchen mit weniger Energie. In jedem Laut sind Grundfrequenz sowie ihre ganzzahligen Vielfachen, die so genannten Harmonischen zu erkennen. Die Abbildung zeigt, dass die Grundfrequenz mit zunehmendem Alter abnimmt und sich geschlechtsspezifische Unterschiede erst bei Heranwachsenden ausbilden (aus: Ey et al. 2007 a).

wie bei einem heranwachsenden Tier (Fischer et al. 2004). Aber nicht nur in der akustischen Struktur, sondern auch im Einsatz unterscheiden sich ranghohe und rangniedere Tiere: Bei den Interaktionen zwischen männlichen Tieren rennen die beteiligten Tiere oft über weite Strecken, verfolgen sich oder springen in Baumkronen umher und rütteln an Ästen. Dieser Verhaltenskomplex wird im Allgemeinen als ein Zeichen der Ausdauer und Konstitution des Tieres angesehen. Hochrangige männliche Tiere äußern mehr Laute pro Serie, ihre Rufserien sind länger, und diese Rufserien werden häufiger vorgetragen als bei rangniedrigen Tieren. Letztere weisen überhaupt eine geringere Neigung auf, sich an solchen Auseinandersetzungen zu beteiligen; aber auch für sie gilt, dass Tiere ähnlichen Ranges generell häufiger Handel eingehen als solche mit weit auseinanderliegenden Rängen (Kitchen et al. 2003).

Einerseits also ist die Grundstruktur der Lautmuster bei Affen angeboren, andererseits wirken neben Alter und Geschlecht zahlreiche

Einflussfaktoren auf die Feinstruktur. Wie soeben erläutert, können dies der Rang bzw. die Kampfkraft eines Tieres sein, sein hormoneller Zustand oder auch seine Motivation oder Erregung. Weitere Studien deckten zudem Variationen in Zusammenhang mit dem Kontext auf, in dem ein Lautmuster eingesetzt wird. Individuelle Unterschiede in Lautmustern sind ein weiterer Grund für die Variabilität der Lautäußerungen bei Primaten. Darüber hinaus spielen die Artikulation, die Zugehörigkeit zu einer bestimmten Gruppe oder Population eine Rolle (Fischer 2003).

Die beschränkte Flexibilität in der vokalen Kommunikation nicht-menschlicher Primaten ist in der neuronalen Kontrolle der Lautäußerungen begründet. Wie Arbeiten von Uwe Jürgens und Kollegen vom Deutschen Primaten-Zentrum (DPZ) in Göttingen zeigten, fehlen Affen direkte Verbindungen vom Cortex zum Kehlkopf (Jürgens 2002). Es ist ihnen daher nicht möglich, die Struktur ihrer Laute willkürlich zu kontrollieren. Vielmehr gibt es eine wichtige Integrationsstelle, das so genannte zentrale Höhlengrau, das Input vom vorderen limbischen Cortex sowie sensorischen und motivationalen Arealen erhält. Das zentrale Höhlengrau sendet seinerseits Projektionen zu Motorkernen im Hirnstamm, die die Atmung, die Artikulation sowie den Stimmeinsatz steuern. Dabei handelt es sich im Wesentlichen um Mustergeneratoren, die mehr oder weniger festgelegte Lautstrukturen produzieren. Die willkürliche Steuerung ist auf den Einsatz der Lautmuster sowie die Lautdauer beschränkt.

Entwicklung der Lautwahrnehmung und Lernen von Lauten

Doch nicht bei allen Tierarten ist die Struktur der Laute derart festgelegt. Singvögel zum Beispiel müssen den arttypischen Gesang lernen. In unseren Breiten hören die Jungvögel nach dem Schlüpfen den Gesang des Vaters oder anderer männlicher Vögel in der Nähe. Die Jungtiere merken sich den Gesang; man spricht davon, dass sie eine Schablone ausbilden. Im Herbst und Winter erfolgt dann eine Phase erhöhter vokaler Aktivität, wobei die Laute zunächst noch recht unstrukturiert sind. Mit zunehmender Praxis bildet sich dann immer deutlicher der arteigene Gesang aus; man spricht davon, dass der Gesang kristallisiert (Marler und Slabbekorn 2004). Delphine sind

anders als Affen in der Lage, die Lautmuster ihrer Artgenossen zu imitieren. Dies wurde am Beispiel der Nachahmung von so genannten Signaturpfeifen gezeigt. Bei Großen Tümmlern besitzt jedes Tier einen für ihn typischen Signaturpfeiff. In einer Studie konnte Vincent Janik von der schottischen Universität St. Andrews zeigen, dass die Tiere sich gewissermaßen gegenseitig rufen konnten: Wenn ein Delphin den eigentlich für ein anderes Gruppenmitglied typischen Pfeiff von sich gab, war die Wahrscheinlichkeit höher, dass beide danach miteinander interagierten (Janik 2000). Und schließlich sollte man sich nochmals vergegenwärtigen, wie sich unsere eigene Sprachfähigkeit mit dem Alter entwickelt. Neugeborene verfügen wie nichtmenschliche Primaten über ein Repertoire von angeborenen Lauten wie Schreien, Jammern, Stöhnen, Gurren etc. Mit zunehmendem Alter kommt dann das Lachen hinzu. Diese Laute unterscheiden sich nicht bei hörenden und gehörlosen Kindern, wie eine Studie von Elisabeth Scheiner (DPZ Göttingen) und Kollegen zeigte (Scheiner und Hammerschmidt 2008). Bei hörenden Kindern kommt dann als Vorstufe der Sprache das kanonische Lallen hinzu, bei dem Silben redupliziert werden (>da-da-da«). Auffällig ist, dass Kinder mit eingeschränkter Hörfähigkeit erst signifikant später anfangen zu lallen als hörende Kinder.

Im Gegensatz zur angeborenen Struktur der Lautmuster sind adäquate Reaktionen auf Lautmuster nicht von Geburt an zu beobachten, sondern entwickeln sich erst im Laufe der Individualentwicklung. Eine der frühesten Untersuchungen hatte die Entwicklung der Reaktionen auf verschiedene Alarmrufe bei Grünen Meerkatzen zum Thema. Grüne Meerkatzen äußern als Reaktion auf die drei wichtigsten Raubfeinde – Leopard, Adler und Python – akustisch unterschiedliche Alarmrufe, die bei Artgenossen jeweils adaptive Reaktionen auslösen. Als Reaktion auf einen Adler-Alarmruf blicken die Tiere zum Beispiel in den Himmel oder laufen in einen Busch, wogegen sie sich als Reaktion auf einen Schlangen-Alarmruf auf die Hinterbeine stellen und den Boden um sich herum nach der Schlange absuchen. Playbackexperimente zeigten, dass diese Rufe allein, auch ohne die Präsenz eines Raubfeindes, die entsprechenden Reaktionen hervorrufen (Seyfarth et al. 1980). Das Alarmrufsystem der Grünen Meerkatze wurde ursprünglich als ein Beispiel für so genannte symbolische oder semantische Kommunikation im Tierreich angesehen. Bald aber wurde darauf hingewiesen, dass nur zwei Komponenten Voraussetzung für ein solches Kommunikationssystem sind: erstens eine gewisse Spezifität in

der Lautgebung und zweitens die Fähigkeit der Empfänger, das Auftreten dieser Signale mit bestimmten Ereignissen zu verknüpfen. Demnach ist es unerheblich, ob der Lautgeber die Intention hat, eine bestimmte Information zu übermitteln, oder ob die Lautäußerung lediglich Ausdruck eines spezifischen emotionalen Zustands ist. Entsprechend wurde der Begriff der funktionalen Referentialität geprägt, um auf den fundamentalen Unterschied der Rolle von Sender und Empfänger in einem derartigen Kommunikationssystem hinzuweisen.

Wie entwickeln sich nun die Reaktionen der Tiere mit dem Alter? Dazu führten Robert Seyfarth und Dorothy Cheney Versuche durch, bei denen sie jungen Grünen Meerkatzen Alarmrufe vorspielten. Im Alter von 3-4 Monaten liefen die Jungtiere in der Regel zu ihrer Mutter, unabhängig davon, welcher Laut präsentiert wurde. Im Alter von 4-5 Monaten zeigten manche der getesteten Tiere als Reaktion auf das Vorspiel von Alarmrufen falsche Reaktionen, zum Beispiel blickten sie nach dem Vorspiel eines Schlangenalarmrufes in den Himmel. Jungtiere in diesem Alter reagierten eher adäquat, wenn sie zuerst ein adultes Tier angeguckt hatten. Im Alter von 6-7 Monaten schließlich verhielten sie sich wie adulte Tiere. Grüne Meerkatzen reagieren auch auf die Alarmrufe anderer Arten, zum Beispiel auf die Rufe von Glanzstaren, die diese als Reaktion auf Bodenfeinde wie etwa Leoparden äußern. Alarmrufe von Staren treten in verschiedenen Habitaten mit unterschiedlicher Häufigkeit auf. Jungtiere, die in einem Habitat leben, in dem die Staren-Alarmrufe häufig vorkommen, reagieren früher auf solche Alarmrufe als Tiere, die diese nur selten hören. Die Ergebnisse dieser Experimente wiesen bereits auf eine gewisse Plastizität in der Entwicklung der Reaktionen bei Primaten hin (Übersicht in Fischer 2002).

In Zusammenarbeit mit Dorothy Cheney und Robert Seyfarth führte ich eine weitere Studie zur Ontogenese der Unterscheidung von Lautmustern durch. Als Modell dienten die Bell-Laute adulter weiblicher Paviane. Akustische Analysen hatten gezeigt, dass es zwei verschiedene typische Varianten gibt, die in verschiedenen Kontexten vorkommen: einerseits die bereits besprochenen harmonischen und lang gezogenen Kontaktrufe und andererseits die schrillen und kürzeren Alarmrufe, wobei jedoch auch intermediäre Varianten häufig vorkommen. Für das Vorspiel wählten wir Lautmuster, die sich gemäß der Analyse akustisch deutlich bzw. akustisch nur wenig unterschieden.

Sobald wir ein geeignetes Tier für das Experiment gefunden hatten, versteckte einer der beiden Beobachter den Lautsprecher in einem rechten Winkel zur Blickrichtung des Testtieres, während der andere Beobachter das Zeichen zum Beginn des Experimentes gab und das Tier mit einer Videokamera filmte. Die Tiere mussten zum Zeitpunkt des Experimentes außerhalb der Sichtweite der Mutter sein. Die Jungtiere zeigten deutlich unterschiedliche Reaktionen auf das Vorspiel typischer Alarm- und Kontaktrufe. Am stärksten reagierten sie auf das Vorspiel von Alarmrufen und am wenigsten auf das Vorspiel von typischen Kontaktrufen – allerdings nur, wenn diese von einem nicht verwandten Weibchen geäußert worden waren. Die Reaktionen auf das Vorspiel der intermediären Varianten löste mittlere Reaktionen aus, wobei nicht zu klären war, ob die Tiere zwischen intermediären Alarm- und intermediären Kontaktrufen unterschieden. Um nun die ontogenetische Entwicklung der Lautmusterunterscheidung zu überprüfen, untersuchten wir in einer Längsschnittstudie, ab welchem Alter die Tiere adäquat auf typische Versionen von Kontakt- und Alarmrufen reagieren. Zu diesem Zweck spielten wir sechs Jungtieren im Alter von zweieinhalb, vier und sechs Monaten wiederholt einen typischen Kontaktruf oder einen typischen harschen Alarmruf vor. Da die Tiere in der jüngsten Altersstufe sich in der Regel in der Nähe ihrer Mutter aufhielten, führten wir alle Experimente durch, wenn die Jungtiere in Sichtweite der Mutter waren, nicht aber in Körperkontakt zu ihr standen. Jungtiere im Alter von zweieinhalb Monaten reagierten weder auf das Vorspiel eines typischen Alarmrufes noch auf das Vorspiel eines typischen Kontaktrufes. Im Alter von vier Monaten blickten einige Jungtiere nach der Präsentation der Laute in die Richtung des Lautsprechers. Allerdings gab es keinen Unterschied in Abhängigkeit davon, welcher Laut präsentiert wurde. Im Alter von sechs Monaten schließlich reagierten die Jungtiere deutlich auf das Vorspiel eines harschen Alarmrufes, wohingegen sie nach dem Vorspiel eines typischen Kontaktrufes keine Reaktion mehr zeigten. In diesem Stadium verhielten sie sich demnach wie adulte Tiere (Fischer et al. 2000). Die untersuchten Pavian-Jungtiere waren also in einem Alter von etwa sechs Monaten in der Lage, zwischen verschiedenen Varianten innerhalb eines Lauttyps zu unterscheiden. Dabei schienen sie verschiedene Phasen zu durchlaufen: Zuerst zeigten sie keinerlei Reaktionen auf Lautmuster, die – wie Alarm- und Kontaktrufe – in der Regel in einiger Entfernung auftraten. Erst im Alter von etwa vier Monaten reagierten

sie dagegen auf Lautmuster, die in einiger Entfernung vorkommen. In einem zweiten Schritt begannen sie dann, verschiedene Varianten innerhalb dieses selben generellen Lauttyps zu unterscheiden. Die vorliegenden Experimente ermöglichten für sich genommen keine Entscheidung der Frage, ob es sich hier um einen Reifungs- oder einen Lernprozess handelt. Allerdings weisen mehrere Studien auf den Einfluss von Erfahrung auf die Reaktion auf verschiedene Lautmuster hin: So lernten zum Beispiel weibliche Rhesus- und Japanmakaken, artfremde »Stiefkinder« an der Stimme zu erkennen. Außerdem achten Primaten auch auf die Alarmrufe von Vögeln, Huftieren oder Vertretern anderer Primatenspezies.

Sowohl die Studie an Grünen Meerkatzen als auch die an Bärenpavianen ergab, dass die Jungtiere etwa in einem Alter von einem halben Jahr adäquate Reaktionen auf verschiedene Lauttypen zeigen. Das heißt, sie waren in diesem Alter in der Lage, diese korrekt zu klassifizieren. Unklar blieb bei diesen Studien, ob die beobachtete Altersgrenze bei diesen beiden Altweltaffenarten ein Ausdruck der allgemeinen Entwicklung ist oder ob sie vielmehr auf die untersuchten Lauttypen zurückzuführen ist. Diese werden im Gegensatz zu anderen Lauttypen in der Regel über weite Distanzen geäußert, so dass es für die Jungtiere möglicherweise schwierig ist, die Ursache des Rufens zu erfassen und dementsprechend die Laute mit Bedeutung zu belegen. Um diese Frage zu klären, initiierte ich eine Studie, in der ich die Entwicklung der Erkennung der mütterlichen Stimme untersuchte. Diese sollte eine der frühesten Erkennungsleistungen sein, die die Jungtiere ausbilden. Zudem konnte ich einen Mangel an Erfahrung ausschließen, da die Tiere in diesem Alter die meiste Zeit mit ihren Müttern verbringen. Ich führte diese Studie an Berberaffen durch, weil bei dieser Art die Jungtiere schon im Alter von wenigen Tagen längere Zeitabschnitte in der Obhut anderer Tiere als der Mutter verbringen. Die Jungtiere werden dabei oft am Fußgelenk gehalten, so dass sie nur einen kleinen Aktionsradius haben. Da die Mutter in solchen Fällen oft nicht zugegen ist, können Experimente an Neugeborenen von einem frühen Alter an durchgeführt werden, ohne in die Gruppe eingreifen zu müssen.

Als Stimuli in den Experimenten wählte ich kurze Schreie. Dieser Lauttyp kommt häufig vor, meist bei Streit in der Gruppe. Die Schreie von Berberaffen sind, ebenso wie die Rufe zahlreicher anderer Primatenarten, individuell distinkt. In den Experimenten präsentierte ich

den Jungtieren entweder Rufe der Mutter oder eines anderen weiblichen Tieres der gleichen Gruppe. Ich führte die Experimente blind durch, das heißt, ich wusste nicht, aus welcher Bedingung die vorgespielten Lautmuster stammten. Die Reaktionen der Tiere wurden mit einer Videokamera aufgezeichnet und später immer noch blind in Bezug auf die experimentelle Bedingung ausgewertet. Neben der Identität des rufenden Tieres betrachtete ich den Einfluss anderer Effekte, wie zum Beispiel die Ausrichtung des Lautsprechers in Bezug auf das Tier, die Distanz zum Lautsprecher sowie die Anzahl der Tiere in der Nähe des Versuchstieres. In die Analyse der Reaktionen der Jungtiere ging als Variable zusätzlich noch das Alter ein.

Um zunächst die Brauchbarkeit des experimentellen Ansatzes zu überprüfen, führte ich eine Reihe von Experimenten an Einjährigen durch. In den Experimenten reagierten die Tiere signifikant stärker auf das Vorspiel der mütterlichen Laute als auf das Vorspiel eines nicht verwandten Weibchens. In zwei der neun Experimente näherte sich das Versuchstier sogar an den versteckten Lautsprecher an; in einem Fall zeigte das Tier dabei freundliche mimische Gesten wie Schmatzen und Zähneklappern. Die Versuchsbedingung allein erklärte dabei die beobachtete Varianz in den Reaktionen am besten, während die übrigen Variablen keinen Einfluss hatten. Der gewählte Versuchsansatz schien also geeignet zu sein, Unterschiede in Reaktionen auf mütterliche Rufe bzw. die Rufe nicht verwandter weiblicher Tiere festzustellen.

Die Jungtiere wurden, ähnlich wie in der oben genannten Studie an Pavianen, in drei Alterskategorien eingeteilt, wobei das Alter im Mittel vier, zehn oder sechzehn Wochen betrug. Im Alter von vier Wochen reagierten die Jungtiere in der Regel nicht auf das Vorspiel von Lauten. Ab dem Alter von zehn Wochen zeigten sie bereits deutlich stärkere Reaktionen nach der Präsentation der mütterlichen Laute als nach dem Vorspiel der Laute anderer Weibchen. Sie schienen also in der Lage zu sein, ihre Mutter an der Stimme zu erkennen und ihre Rufe von den Rufen anderer weiblicher Tiere zu unterscheiden. Wie auch in den Experimenten mit den Bärenpavianen zeigten die Tiere in der jüngsten Alterskategorie keine Reaktionen auf das Vorspiel der Lautmuster (Fischer 2004).

Die Produktion von Lauten und ihr Verständnis entwickeln sich also entlang sehr unterschiedlicher ontogenetischer Linien. Neugeborene Berberaffen produzieren vom ersten Tag ihres Lebens an eine Vielzahl von verschiedenen Lautmustern, die sich nur unwesentlich

von den Lauttypen unterscheiden, die von älteren Tieren geäußert werden. Im Gegensatz dazu setzt die Erkennung von Lautmustern erst später ein. Offensichtlich gibt es Unterschiede in der Geschwindigkeit der Entwicklung der Reaktionen in Abhängigkeit vom Lauttyp. Während junge Berberaffen in einem Alter von zehn Wochen die Stimme ihrer Mutter erkennen, reagieren Pavian-Jungtiere zu diesem Zeitpunkt noch überhaupt nicht auf das Vorspiel von Kontakt- und Alarmrufen. Beide Arten weisen im Übrigen aber eine ähnliche Frühentwicklung auf. Diese Ergebnisse weisen auf eine gewisse Plastizität in der Entwicklung der Unterscheidung von Lautkategorien hin, die vermutlich vom Umfang der Erfahrung sowie der Bedeutsamkeit der Stimuli abhängt.

Verschiedene Komponenten sind Voraussetzung für die Entwicklung einer adäquaten Reaktion. Zum einen müssen die Jungtiere in der Lage sein, die Lautmuster von anderen Geräuschen zu unterscheiden und die Richtung der Schallquelle zu lokalisieren. Hier spielen vermutlich Reifungsprozesse eine wichtige Rolle, wobei dies im Detail nicht untersucht ist. In einem nächsten Schritt müssen sie ihre Aufmerksamkeit auf die arteigenen Lautmuster richten, wobei soziales Lernen dabei wahrscheinlich eine wichtige Rolle spielt. Innerhalb eines assoziativen Lernprozesses müssen die Jungtiere schließlich eine spezifische Lautäußerung mit dem Individuum verbinden, welches den Laut äußert, bzw. mit dem Kontext, in dem diese Lautäußerung auftritt. Selbstverständlich können sie die adäquaten Reaktionen erst dann produzieren, wenn sie motorisch dazu in der Lage sind.

Der Mangel an Reaktionen in der jüngsten Altersklasse, den ich sowohl bei Berberaffen als auch bei Bärenpavianen beobachten konnte, ist vermutlich unabhängig vom Ruftyp. Allerdings kann das Fehlen einer Reaktion nicht einfach einer motorischen Unfähigkeit zugeschrieben werden: Neugeborene blicken bereits im Alter von einer Woche sich bewegenden Objekten in ihrer Umgebung hinterher. Zudem reagieren sie bereits in der ersten Lebenswoche auf soziale Signale wie Anschmatzen mit entsprechenden mimischen Gesten wie langsamen Kaubewegungen. Motorische Beschränkungen können also ausgeschlossen werden. Auch wenn junge Berberaffen schon im Alter von etwa zehn Wochen verstärkt auf die mütterliche Stimme reagierten, so ist dies doch sehr viel später als zum Beispiel bei Pelzrobben. Junge Pelzrobben lernen bereits in der ersten Lebenswoche, die Stimme ihrer Mutter von der anderer weiblicher Tiere zu unterscheiden, noch bevor

die Mutter ihren ersten ausgedehnten Fischzug antritt. In diesem Sozialsystem beruht die Kontaktaufnahme auf gegenseitigem Rufen, welches der Lokalisierung in der Kolonie dient. Vermutlich liegt im Sozialsystem der Robben ein hoher selektiver Druck in Richtung einer sehr frühen Erkennung und Reaktion. Bei Primaten hingegen ist in den ersten Lebenswochen die Mutter für die Wiederherstellung des Kontaktes verantwortlich. Erst wenn die Jungtiere einen größeren Bewegungsradius haben und nicht konstant durch adulte Tiere betreut werden, zeigen sich differentielle Reaktionen auf die Rufe der Mutter.

Inwieweit der Lernprozess durch soziale Erfahrung bestimmt wird, bedarf weiterer Beobachtungen. Die vorliegende Datenlage ist zum Teil widersprüchlich: Einerseits stehen die Reaktionen der Berberaffen-Jungtiere nicht in Zusammenhang mit etwaigen Reaktionen des Betreuers. Andererseits zeigen junge Grüne Meerkatzen eher adäquate Reaktionen auf verschiedene Alarmrufe, wenn sie zuvor ein adultes Tier angeblickt haben (Fischer 2002). Junge Paviane reagieren wiederum im Gegensatz zu adulten Tieren auf intermediäre Alarm- und Kontaktrufe – sie kopieren also nicht das Verhalten der Adulten. Auch junge Berberaffen zeigen in der Regel stärkere Reaktionen auf Alarmrufe als Adulte. Möglicherweise neigen Jungtiere, sobald sie einmal gelernt haben, auf bestimmte Reize zu reagieren, eher dazu, diese zu generalisieren. Die Evidenz für jede Form von »pädagogischem Verhalten« oder der aktiven Vermittlung von Wissen bleibt mager (Fischer 2007). Ebenso ungeklärt bleibt bislang, wie hoch der Einfluss sozialer Erfahrung ist. Nichtmenschliche Primaten folgen regelmäßig den Blicken anderer Gruppenmitglieder, aber auch denen menschlicher Experimentatoren. So folgen Rhesusaffen ab einem Alter von fünf-einhalb Monaten den Blicken menschlicher Experimentatoren. Möglicherweise setzt dieses so genannte *gaze following* zwischen Artgenossen aber bereits in einem früheren Alter ein.

Zusammenfassend lässt sich also feststellen, dass junge Primaten lernen müssen, die Lautmuster ihrer Artgenossen mit Bedeutung zu belegen. Dabei hängt die Entwicklung der Reaktionen anscheinend davon ab, welcher Art die Lautmuster sind, in welchen Kontexten sie auftreten und wie häufig sie vorkommen, wobei quantitative Daten zu diesen Zusammenhängen noch fehlen. Unklar bleibt bislang, ob junge Primaten eine angeborene Disposition besitzen, sich arteigenen Lautmustern zuzuwenden und inwiefern soziale Erfahrung den Lernprozess beeinflusst. Leider lassen sich diese Fragen im Freiland nur schwer

oder gar nicht untersuchen, da man kaum Kontrolle darüber hat, welche Erfahrungen ein Tier in welchem Alter gemacht hat und wie sich dies auf das Verständnis von Lautmustern auswirkt.

Wortlernen beim Haushund

Wer sich – wie ich – davon beeindrucken lässt, dass Grüne Meerkatzen drei verschiedene Alarmrufe unterscheiden können, musste hellhörig werden, als in der Presse von einem Hund die Rede war, der angeblich 70 verschiedene Spielzeuge beim Namen kannte. Dieser Hund hatte in der Sendung »Wetten, dass...?« einen großen Auftritt, über den auch in der Presse berichtet wurde. Ich nahm daraufhin Kontakt zu den Besitzern von Rico auf, die uns herzlich zu sich nach Hause einluden, um Rico kennen zu lernen. Zusammen mit Juliane Kaminski, die damals auch am Max-Planck-Institut für evolutionäre Anthropologie forschte, fuhr ich nach Dortmund, um Ricos Fähigkeiten zu begutachten. Als erstes galt es zu klären, ob Rico die Spielzeuge wirklich beim Namen kennt oder ob er – ähnlich wie das berühmte Pferd »Der Kluge Hans« – auf unmerkliche Zeichen seiner Besitzerin achtet. Der Kluge Hans war ein Pferd des Lehrers Wilhelm von Osten, das angeblich rechnen konnte. Das Pferd beantwortete Aufgaben mit dem Klopfen eines Hufes oder durch Nicken bzw. Schütteln seines Kopfes. Anfang des letzten Jahrhunderts erlangten der Kluge Hans und seine Fähigkeiten erstaunliche Berühmtheit. Damals wurde eigens eine Kommission der Berliner Friedrich-Wilhelms-Universität eingesetzt, die zu dem Schluss kam, dass der Kluge Hans tatsächlich rechnen konnte. Nur der Psychologe Oskar Pfungst ließ sich nicht so recht überzeugen und stellte dem Pferd eine Aufgabe, deren Antwort Herr von Osten selbst nicht kannte. Es zeigte sich, dass das Pferd einfach weiter mit dem Huf scharrte. Die vermeintlichen Rechenkünste des Klugen Hans waren also eher Ausdruck der genauen Beobachtungsgabe des Pferdes, das gelernt hatte auf unwillkürliche (und unwissentlich gegebene) Zeichen seines Herrn zu achten.

Die Experimente, die Juliane Kaminski und ich zusammen mit Josep Call durchführten, dienten zunächst der Klärung, ob Rico nicht vielleicht doch ein zweiter Kluger Hans ist (Kaminski et al. 2004). Als wir die Studie begannen, besaß Rico bereits über 200 verschiedene Spielzeuge, zumeist Stofftiere, aber auch verschiedene Bälle und einen

Gummi-Hamburger. Die 200 Spielzeuge wurden in 20 Sets zu je zehn Objekten aufgeteilt. Jeweils zehn Spielzeuge wurden in einem Raum verstreut, es wurde eine Videokamera aufgebaut, und dann begann das typische Suchspiel, das Ricos Besitzerin, Suse Baus, immer mit ihm spielte. Von einem anderen Raum aus, von dem man nicht in das Zimmer blicken konnte, in dem die Spielzeuge lagen, forderte Suse Baus Rico auf, eines der Spielzeuge zu holen. Die ganze Prozedur ist recht ritualisiert. »Riiiiicoo!«, begann Suse jede Runde, »wo ist denn der Weihnachtsmann?« Dann trabte Rico los, und in der Regel kam er auch mit dem Weihnachtsmann im Maul zurück. Insgesamt führten wir 40 solche Experimente durch, und in 37 Fällen holte Rico das richtige Spielzeug. Wohlgemerkt, es war niemand im Raum, als er das Spielzeug aufnahm. Aus diesen Experimenten schlossen wir, dass Rico tatsächlich den Namen bzw. das Wort mit dem jeweiligen Spielzeug verbindet (Kaminski et al. 2004).

Es drängte sich damit die Frage auf, wie er die Verbindung herstellt. Unsystematische Beobachtungen zeigten, dass er sehr schnell und ohne großes Training ein neues Wort mit einem neuen Gegenstand assoziiert. So brachten wir Rico am ersten Tag einen kleinen Plüschpapagei mit, den Juliane ihm mit den Worten gab, dass dies der »Josep« sei. Er spielte ein bisschen mit dem Papagei und fortan holte er dieses Stofftier, wenn man ihn aufforderte, den »Josep« zu bringen. Interessanterweise brachte er genau dieses Tier Juliane, als sie Wochen später wieder bei der Familie Baus zu Besuch war, ohne dass sie ihn dazu aufgefordert hatte, und Suse Baus bestätigte uns, dass Rico sich nicht nur daran erinnerte, wie ein Spielzeug heißt, sondern auch, von wem er es bekommen hatte. Das haben wir aber nie systematisch untersucht. Was uns mehr interessierte, war die Frage, ob Rico eine bestimmte Form des Wortlernens beherrscht, die bislang nur dem Spracherwerb bei Kindern zugeordnet wurde. Es handelt sich dabei um das Schnelle Zuordnen oder *fast mapping*. Die ersten Experimente dazu wurden von Susan Carey und Elsa Bartlett Ende der 1970er Jahre in einem Kindergarten durchgeführt (Carey und Bartlett 1978). Ein Experiment begann damit, dass ein rotes und ein olivgrünes Tablett auf ein kleines Tischchen gelegt und ein Kind gebeten wurde, das »chromfarbene Tablett zu holen, nicht das rote«. Und normalerweise brachten die Kinder dann das olivgrüne Tablett, denn sie wussten ja schon, dass die eine Farbe Rot heißt, also der andersfarbige Gegenstand gemeint sein muss und im Übrigen der Name dieser anderen Farbe wohl »chrom-

farben« ist. Wenn man sie hinterher fragte, welches denn das chromfarbene Tablett sei, dann zeigten sie auf das olivgrüne; man kann ihnen aber auch einen ganz anderen Gegenstand in der gleichen Farbe geben und sie fragen, wie die Farbe heißt. Kinder in einem Alter von etwa drei Jahren werden dann in der Regel »chromfarben« sagen; sie sind also in der Lage, durch dieses Ausschlussverfahren einen neuen Namen einer noch unbenannten Farbe zuzuordnen und sich diesen Namen dann auch zu merken. Frühe Debatten rankten sich um die Frage, ob das Schnelle Zuordnen wohl nur auf den Bereich von Farben beschränkt sei oder ein allgemeines Prinzip beim Zuordnen von Namen zu Dingen ist; letzteres wurde später durch weitere Studien bestätigt.

Uns interessierte nun, ob Rico auch zum schnellen Zuordnen in der Lage ist. Dazu führten wir ein Experiment durch, dessen Design gegenüber den Experimenten mit Kindern natürlich etwas abgewandelt werden musste. Schließlich konnten wir Rico nicht fragen, wie das neue Spielzeug heißt. Wieder nutzten wir Ricos Lust am Suchspiel. Zunächst wurden in einem Zimmer verschiedene Spielzeuge versteckt, die er schon kannte. Dazu wurde dann ein neues Spielzeug gelegt. Genau wie im ersten Experiment wurde auch hier Ricos Verhalten auf Video aufgezeichnet, und niemand war im Raum, der ihn in irgendeiner Weise beeinflussen konnte. Zunächst forderte Suse Baus Rico auf, ein oder zwei Spielzeuge zu holen, die er schon kannte. Und dann bat sie ihn, das neue zu holen. Sie fragte zum Beispiel »Wo ist denn der Lessing?« In sieben von zehn Fällen brachte er dann auch das jeweils neue Spielzeug, wenn seine Besitzerin ein neues Wort nannte. Auch wenn sich dies zunächst beeindruckend anhört: Es gibt eine Reihe vergleichbarer Lernexperimente, zum Beispiel an Seelöwen, die ein solches Lernen durch Ausschluss demonstrieren konnten. Die Forschungen zum Lernen bei Tieren und zum Spracherwerb bei Kindern sind teilweise parallel verlaufen, ohne dass gleich erkannt wurde, dass man sich eigentlich mit dem gleichen Lernprinzip beschäftigte.

Was uns daher viel mehr interessierte, war die Frage, ob Rico sich auch den Namen des Spielzeugs gemerkt hatte, das er im Ausschlussverfahren in einem einzigen Experiment identifiziert hatte. Dazu wurden die von Rico richtig zugeordneten Spielzeuge für vier Wochen in einem Schrank versteckt. In einer weiteren experimentellen Sitzung testeten wir nun, ob er sich an den Namen erinnerte. Um das zu prüfen, legten wir vier neue und vier alte Gegenstände und (zum Beispiel) den »Lessing« in einen Raum. Erst wurde Rico wieder aufgefordert,

ein oder zwei gut bekannte Spielzeuge zu holen, und dann bat seine Besitzerin ihn, den »Lessing« zu holen. In sechs derartigen Experimenten brachte Rico dreimal das richtige Spielzeug – und das nach vier Wochen! Dieses Ergebnis ist zwar statistisch nur marginal signifikant, aber wir waren dennoch ziemlich beeindruckt.

Bei einer Wiederholung der Experimente, bei der am gleichen Tag abgefragt wurde, ob er sich den Namen merken konnte, war dies in vier von sechs Runden der Fall. Rico hatte sich also gemerkt, dass ein Wort, welches seine Besitzerin nur ein einziges Mal genannt hatte, zu einem Spielzeug gehörte, das er durch Ausschluss identifiziert hatte. Wurde an demselben Tage, an dem ein neuer Gegenstand durch Ausschlussverfahren erstmals identifiziert wurde, das Experiment wiederholt, so zeigte sich, dass Rico in vier von sechs Runden die richtige Wahl traf.

Auch wenn er sich nicht an alle Namen erinnern konnte, so muss man bedenken, dass Kinder sich auch etwa nur die Hälfte der Namen merken und Erwachsene noch viel schlechter abschneiden. In jedem Fall ist das Schnelle Zuordnen nicht auf den Spracherwerb bei Kindern beschränkt. Wir führen diese Fähigkeit nicht auf ein besonderes Lernmodul zurück, sondern auf die Kombination einfacher Mechanismen. Voraussetzungen für das Schnelle Zuordnen wären demnach eine Erwartungshaltung, dass Objekte Namen haben, ferner die Fähigkeit, eine Auswahl durch Ausschluss vorzunehmen, und schließlich die Möglichkeit, sich die Zuordnung zu merken.

Studien zu syntaktischen Fähigkeiten bei Tieren

Wir haben bislang gesehen, dass die Bedeutung von Lauten diesen im Wesentlichen durch die Empfänger zugeschrieben wird. Wie sieht es nun mit den syntaktischen Fähigkeiten der Tiere aus? Können sie die ihnen zur Verfügung stehenden Lautelemente kombinieren und damit unterschiedliche Botschaften übermitteln? Bei der Beantwortung dieser Frage muss man sowohl die Sender- als auch die Empfängerseite betrachten. Einiges Aufsehen erregte eine Serie von Studien zu der Kommunikation von Weißnasenmeerkatzen, die Klaus Zuberbühler und Kollegen (*University of St. Andrews*) durchführten (Arnold und Zuberbühler 2008). Diese Meerkatzen verfügen über zwei verschiedene Ruftypen, *backs* und *pyows*, die sie in verschiedenen Kontexten

einsetzen. Zuberbühler und Kollegen beobachteten, dass die Tiere nach dem Erscheinen eines Adlers vornehmlich (aber nicht ausschließlich) Rufserien äußerten, die mit *backs* begannen und manchmal von *pyows* gefolgt wurden. Dagegen riefen die Tiere vornehmlich *pyow*-Serien, die in *backs* übergingen, wenn sie einen Leopard gesehen hatten. Kompliziert wird die ganze Angelegenheit noch dadurch, dass die männlichen Tieren vorwiegend gemischte Serien vortrugen, wenn sie eine Gruppenbewegung initiieren wollten. Es bleibt allerdings die Frage, ob ein solches System mit der menschlichen Grammatik vergleichbar ist.

Betrachtet man die statistische Verteilung der *pyows* und *backs* in den verschiedenen Sequenzen, so drängt sich eher der Schluss auf, dass es sich hier um ein probabilistisches System handelt, bei dem sich keine eindeutige Regelmäßigkeit erkennen lässt. Im Übrigen sind derartige kontextspezifische Rufkombinationen keine Seltenheit: So geben männliche Bärenpaviane vor einer Auseinandersetzung mit Artgenossen oft typische *roar-grunts* von sich, bevor sie in eine ganze Serie von *wa-hoos* ausbrechen (Hamilton und Bulger 1992). Wie eingangs erwähnt, kommen diese *wa-hoos* aber auch in anderen Kontexten vor, z.B. wenn die Tiere den Kontakt zur Gruppe verloren haben. Wenn ein Zuhörer aber zunächst einen *roar-grunt* hört, weiß er oder sie mit hoher Gewissheit, dass es sich diesmal um eine Auseinandersetzung unter Männchen handelt.

So viel zur Erzeugen grammatikalischer Strukturen. Wie sieht es nun mit dem Verständnis von Signalkombinationen aus? Ein aufsehen-erregender Versuch wurde dazu von einer Gruppe von Forschern um Timothy Gentner (*University of California*, San Diego) an Staren durchgeführt (Gentner et al. 2006). Gentner und Kollegen erzeugten Sequenzen aus verschiedenen Starengesangssilben, denen entweder eine *Finite State Grammar* mit der Formel $(AB)^n$ oder eine kontextfreie Grammatik entsprechend der Formel A^nB^n zugrunde lag. Die erste Grammatik arbeitet mit einem »starren« Element (AB) , das sich zu $ABAB$, $ABABAB$ usw. aneinanderreihen lässt, während sich mit der zweiten Grammatik Sequenzen erzeugen lassen, die formal einer Einbettung einer Struktur in sich selbst entspricht. Dadurch ergeben sich eingeschachtelte Konstruktionen: Die Elementfolge AB wird in AB eingesetzt, wodurch die Folge $AABB$ entsteht. Führt man ein weiteres Mal AB in diese Struktur ein, so ergibt sich $AAABBB$ usw. A und B entsprachen jeweils verschiedenen Silbenexemplaren.

Die Tiere wurden nun trainiert, verschiedene Sequenzen ABAB bzw. AABB in Dressurversuchen zu unterscheiden. Bemerkenswert an dieser Arbeit ist zunächst einmal, wie schwer sich die Tiere mit dieser Aufgabe taten: Die besten vier Tiere brauchten mehr als 20 000 Versuche, bis sie die Aufgaben zuverlässig meisterten – und dies war, bevor sie mit unbekannten Serien getestet wurden. Andererseits könnte man auch sagen, dass immerhin vier von elf Tieren in der Lage waren, die Aufgabe zu meistern. Dazu muss man aber nicht unbedingt ein Verständnis für die grammatikalische Struktur annehmen. Theoretisch reicht es aus, vor allem auf die ersten beiden Silben zu achten und die Anzahl der Silben zu zählen. Noch aufschlussreicher sind Studien der Gruppe um Angela Friederici vom Max-Planck-Institut für kognitive Neurowissenschaften in Leipzig. Sie trainierte menschliche Probanden mit entsprechenden Silbensequenzen (Bahlmann et al. 2008). Die Probanden brauchten dazu im Durchschnitt 35 Minuten. Dies allein ist schon ein gewichtiger Hinweis darauf, wie leicht es uns Menschen fällt, bestimmte Regeln aus Sequenzen zu extrahieren.

An dieser Stelle sei auch ein Rückgriff auf die Sprachversuche mit Menschenaffen erlaubt. Auch hier lohnt sich ein genauer Blick auf die Daten. In die Lehrbücher oder in die Medien schafft es ja oft nur die Kurzform, zum Beispiel die Aussage, die Schimpansin Lana, die von Duane und Sue Rumbeaugh (*Georgia State University*, Atlanta) untersucht wurde, habe spontan die Symbolkombination »Apfel-welcher-ist-orange« getippt, als ihr eine Orange präsentiert wurde. Wenn man allerdings den genauen Dialog verfolgt, der sich vor dieser spezifischen Aussage abspielte, entwickeln sich doch erhebliche Zweifel an der Spontaneität der Äußerung von Lana. Vielmehr wurde Lana so lange korrigiert, bis die schließlich veröffentlichte Zeichenfolge gezeigt wurde. Zudem war es der Trainer, der die Konversation auf das Thema Farbe brachte.

Die derzeit vorliegenden Ergebnisse zur Entwicklung der kommunikativen Fähigkeit bei nichtmenschlichen Primaten legen also den Schluss nahe, dass die Unterschiede in der akustischen Kommunikation zwischen Menschen und diesen Tieren weitaus größer sind als die Gemeinsamkeiten. Die Empfänger sind zwar in der Lage, allen möglichen Lautmustern Bedeutung zuzuschreiben und auch feine Unterschiede in der Lautstruktur wahrzunehmen. Dagegen umfassen die Vokalisationen selbst von Menschenaffen in der Regel nur eine limitierte Anzahl verschiedener genetisch weitgehend festgelegter Laut-

typen. Wegen der mangelnden Flexibilität in der Lautgebung konnte sich auch kein willkürlicher Zusammenhang zwischen Lauten und Kontexten entwickeln, in denen sie eingesetzt werden, oder den Ereignissen, die sie auslösen. Ebenso fehlen überzeugende Belege, dass Affen ihre Lautmuster willkürlich regelhaft einsetzen und damit neue Bedeutung erzeugen.

Das hier Gesagte gilt freilich nicht für alle Tiergruppen, sondern vornehmlich für landlebende Säugetiere. Bei verschiedenen anderen Tiergruppen ist die Fähigkeit zur Imitation wohldokumentiert. Singvögel lernen ihren Gesang vom Vater oder Nachbarn, Delphine können die Pfiffe ihrer Artgenossen nachahmen und es gibt sogar eine schöne Anekdote über eine Robbe, die fluchen konnte (Ralls et al. 1985).

Die menschliche Sprachentwicklung und das FOXP2-Gen

Wie der Mensch zu seiner flexiblen Lautgebung, der ausgeprägten Imitationsfähigkeit und der Fähigkeit, ganze Sätze schon im Voraus zu planen, gekommen ist, liegt nach wie vor im Dunkeln. Die Evolution der menschlichen Sprachfähigkeit verdankt sich offensichtlich einer Vielzahl begünstigender Faktoren. Ein Gen, das womöglich eine entscheidende Rolle für die Entstehung der Sprachfähigkeit spielte und auch heute unabdingbare Voraussetzung ist, ist das so genannte FOXP2-Gen (Vargha-Khadem et al. 1998). Die Bedeutung dieses Gens für die normale Sprachentwicklung beim Menschen wurde bei Untersuchungen der so genannten KE-Familie erkannt, in der über mehrere Generationen hinweg Sprachstörungen diagnostiziert wurden. Es stellt sich heraus, dass in ihr eine bestimmte Punktmutation auftritt, die mit Schwierigkeiten bei der Sprachentwicklung, der Artikulation und der Entwicklung des Sprachverständnisses korreliert. Für eine normale Sprachentwicklung sind zwei intakte Kopien des FOXP2-Gens Voraussetzung.

Wolfgang Enard, Svante Pääbo und Kollegen vom Max-Planck-Institut für evolutionäre Anthropologie verglichen die Sequenzen des FOXP2-Gens bei Menschen, Menschenaffen, Rhesusaffen und Mäusen. Bei Mensch und Maus unterscheiden sich beim FOXP2-Protein lediglich drei der insgesamt 715 Aminosäuren. Der Vergleich der

menschlichen Sequenz mit der von Schimpansen ergab, dass zwei dieser drei Unterschiede erst beim Menschen auftreten und demnach innerhalb der letzten fünf bis sechs Millionen Jahren entstanden sein müssen. Wie eine Analyse direkt benachbarter, nicht für das FOXP2-Protein codierender Bereiche bei heute lebenden Menschen zeigt, war eine der beiden Änderungen sehr wahrscheinlich von selektivem Vorteil während der menschlichen Evolution. Nach Computersimulationen und anderen Untersuchungen dürfte die vorteilhafte Änderung vor etwa 350 000 Jahren erfolgt sein.

Da FOXP2 das erste einzelne Gen ist, welches spezifisch mit der menschlichen Sprachfähigkeit verknüpft ist und auch das erste Gen ist, für das menschenpezifisch eine positive Selektion gezeigt werden konnte, liegt die Hypothese nahe, dass es tatsächlich im Hinblick auf seinen Beitrag zur Sprachfähigkeit selektiert wurde. Die derzeit einzige Möglichkeit, diese Hypothese zu überprüfen, ist, den Effekt der beiden menschenpezifischen Aminosäureänderungen in einem genetisch manipulierbaren System wie der Maus zu studieren. Zu diesem Zweck kann das endogene FOXP2-Gen der Maus an den beiden entsprechenden Stellen durch die menschliche Variante ersetzt werden, so dass sich der Effekt genau dieser Aminosäureänderungen in einem lebenden Organismus untersuchen lässt. Die modifizierten Mausstämmen lassen sich dann mit genetisch identischen, aber nichtmodifizierten Mäusen sowie mit heterozygoten Nachkommen vergleichen. Kurt Hamerschmidt aus meiner Arbeitsgruppe Kognitive Ethologie am Deutschen Primatenzentrum, Wolfgang Enard und ich analysieren im Rahmen einer groß angelegten Studie derzeit den Einfluss der humanen FOXP2-Variante auf die Lautäußerungen von Mäusen (Enard et al. 2009). Wir konnten kürzlich bestätigen, dass FOXP2 u.a. spezifische Wirkungen auf das Gehirn und neuronale Leistungen hat (Veränderungen der Basalganglien, Erhöhung der synaptischen Plastizität) und sogar die Tonhöhe der Ultraschalllaute der Mäuse veränderte, während sich bei über 200 untersuchten physiologischen und morphologischen Parametern kein Einfluss des veränderten FOXP2 ausmachen ließ. Wie nicht anders zu erwarten, können die auf diese Weise »humanisierten« Mäuse also nicht sprechen, doch erhärtet sich, dass mit FOXP2 ein für die Sprachfähigkeit wichtiges Gen gefunden worden ist.

Abschließend möchte ich betonen, dass ich aufgrund der mangelnden Sprachfähigkeit den Tieren nicht attestieren würde, sie seien dumm. Das Erstaunliche ist vielmehr, dass die Tiere sehr viel über ihre

Welt wissen und sich flexibel und intelligent verhalten können, dies aber nicht unmittelbar mit ihren kommunikativen Fähigkeiten verknüpft ist.

Ausblick

Und wie geht die Forschung bei uns weiter? Eines der wichtigsten Projekte, die wir angestoßen haben, ist die Erforschung der Sozialstruktur bei Guineapavianen. Diese leben in Westafrika in einer sehr komplexen sozialen Organisation, wobei sich eine Gemeinschaft von etwa 300 Tieren je nach Saison in mehrere kleinere Gruppen aufspaltet. Dabei ist das System anscheinend sehr flexibel, die Tiere wechseln zwischen verschiedenen Untergruppen hin und her. Was uns interessiert, ist, wie die Tiere durch Kommunikation ihr Gruppenleben regeln. Wir hoffen, dass wir mit diesem Projekt die Hypothese testen können, dass die Kommunikationsfähigkeit ursächlich mit der sozialen Komplexität zusammenhängt.

Mein Dank gilt allen Kooperationspartnern, mit denen ich über die Jahre auf so produktive und erfreuliche Weise zusammenarbeiten durfte. Ganz besonders danken möchte ich darüber hinaus Dorothy Cheney und Robert Seyfarth für die Möglichkeit, in Botswana arbeiten zu dürfen, sowie Kurt Hammerschmidt, mit dem ich über lange Jahre die Freuden und Leiden der akustischen Feinanalyse sowie vieles mehr geteilt habe. Ferner gilt mein Dank allen Mitgliedern der Abteilung Kognitive Ethologie für ihr Engagement und ihren Enthusiasmus. Die hier vorgestellten Arbeiten wurden von der Deutschen Forschungsgemeinschaft (Fi707-2, Fi707-4, Fi707-5, Fi707-7), dem BMBF, dem DAAD und der Leibniz-Gemeinschaft gefördert.

Literatur

- Arnold K., Zuberbühler K. (2008) Meaningful call combinations in a non-human primate. *Current Biology* 18, R202-R203.
- Bahlmann J., Schubotz R.I., Friederici A.D. (2008) Hierarchical artificial grammar processing engages Broca's area. *Neuroimage* 42, S. 525-534.
- Carey S., Bartlett E. (1978) Acquiring a single new word. *Child Language Development*, 15, S. 29.
- Enard W., Gehre S., Hammerschmidt K., Holter S.M., Blass T., et al. (2009) A Humanized Version of Foxp2 Affects Cortico-Basal Ganglia Circuits in Mice. *Cell* 137, S. 961-971.
- Ey E., Hammerschmidt K., Seyfarth R.M., Fischer J. (2007a) Age- and sex-related variations in clear calls of Chacma baboons (*Papio hamadryas ursinus*). *International Journal of Primatology* 28, S. 947-960.
- Ey E., Pfefferle D., Fischer J. (2007b) Do age- and sex-related variations reliably reflect body size in non-human primate vocalizations – a Review. *Primates* 48, 253-267.
- Fischer J., Cheney D.L., Seyfarth R.M. (2000) Development of infant baboons' responses to graded bark variant. *Proceedings of the Royal Society of London Series B – Biological Sciences* 267, S. 2317-2321.
- Fischer J. (2002) Developmental modifications in the vocal behaviour of nonhuman primates. In: *Primate Audition*, hrsg. v. A.A. Ghazanfar. Boca Raton (CRC Press), S. 109-125.
- Fischer J. (2003) *Vokale Kommunikation bei nichtmenschlichen Primaten: Einsichten in die Ursprünge der menschlichen Sprache?* Habilitationsschrift Universität Leipzig, Leipzig.
- Fischer J., Kitchen D.M., Seyfarth R.M., Cheney, D.L. (2004) Baboon loud calls advertise male quality: acoustic features and their relation to rank, age, and exhaustion. *Behavioral Ecology and Sociobiology* 56, S. 140-148.
- Fischer J. (2007) Transmission of acquired information in nonhuman primates. In: *Encyclopedia of Learning and Memory*, hrsg. v. D. Sweatt, R. Menzel, H. Eichenbaum und H. Roediger. Oxford (Elsevier).
- Gentner T.Q., Fenn K.M., Margoliash D., Nusbaum H.C. (2006) Recursive syntactic pattern learning by songbirds. *Nature* 440, S. 1204-1207.
- Hamilton I.W., Bulger J. (1992) Facultative expression of behavioral differences between one-male and multimale savanna baboon groups. *Journal of Primatology* 28, S. 61.
- Herder J.G. (1772) *Abhandlungen über den Ursprung der Sprache*. Leipzig (Reclam).
- Janik V.M. (2000) Whistle matching in wild bottlenose dolphins (*Tursiops truncatus*). *Science* 289, S. 1355.
- Jürgens U. (2002) Neural pathways underlying vocal control. *Neuroscience and Biobehavioral Reviews* 26, S. 235-258.
- Kaminski J., Call J., Fischer J. (2004) Word learning in a domestic dog: Evidence for 'fast mapping'. *Science* 304, S. 1682-1683.

- Kitchen D., Seyfarth R., Fischer J., Cheney D. (2003) Loud calls as indicators of dominance in male baboons (*Papio cynocephalus ursinus*). Behavioral Ecology and Sociobiology 53, S. 374-384.
- Kuckenbug M. (2004) Wer sprach das erste Wort? Stuttgart (Theiss).
- Marler P., Slabbekorn H. (2004) Nature's Music: The Science of Birdsong. San Diego (Elsevier).
- Ralls K., Fiorelli P., Gish S. (1985) Vocalizations and vocal mimicry in captive harbor seals, *Phoca vitulina*. Canadian Journal of Zoology 63, S. 1050-1056.
- Scheiner E., Hammerschmidt K. (2008) The role of auditory feedback in nonverbal vocal behavior of normally hearing and hearing-impaired infants. In: Emotions in the human voice, hrsg. v. K. Izdebski. San Diego (Plural Publishing).
- Seyfarth R.M., Cheney D.L., Marler P. (1980) Monkey responses to three different alarm calls: Evidence of predator classification and semantic communication. Science 210, S. 801-803.
- Vargha-Khadem F., Watkins K.E., Price C.J., Ashburner J., Alcock K.J., Connelly A., Frackowiak R.S.J., Friston K.J., Pembrey M.E., Mishkin M., Gadian D.G., Passingham R.E. (1998) Neural basis of an inherited speech and language disorder. Proceedings of the National Academy of Sciences of the United States of America 95, S. 12 695-12 700.
- Wallmann J. (1992) Aping Language. Cambridge (Cambridge University Press).

A. Julia Stähler

Unsterbliche Elektronen¹

Während der perfekte Urlaub wie im Fluge vergeht und die zwei herbeigesehnten Wochen wie Tage erscheinen, erlebt der Patient den zweiwöchigen Krankenhausaufenthalt wie mehrere Jahre. Offenbar hängt die Wahrnehmung von Zeit, oder vielmehr vom Verstreichen der Zeit, erheblich von den äußeren Bedingungen ab. Besonders schwierig wird es jedoch, sich Zeitspannen vorzustellen, für die es keinen Vergleich aus dem Alltag gibt. Lassen sich Jahrhunderte noch relativ gut in Generationen messen, ist schon der Rückblick über zehntausende Jahre zum Neandertaler schwer zu erfassen. Spätestens das hohe Alter der Erde ist mit 4,5 Milliarden Jahren kaum vorstellbar. Ganz ähnlich verhält es sich auch mit den Zeitskalen, die auf molekularer oder atomarer Ebene relevant sind. So laufen etwa Ladungstransferprozesse, auf denen die Funktionsweise von Solarzellen beruht, innerhalb einiger Femtosekunden (= 0,000 000 000 000 001 Sekunden) ab. Hierbei werden Ladungen, beispielsweise Elektronen, vom Ort A an den Ort B bewegt (transferiert). Das Verständnis vom Ladungstransfer spielt auch für die Fortentwicklung von Computerchips eine große Rolle: Die Ausmaße ihrer Kernbauteile, der Transistoren, konnten in der Vergangenheit stetig reduziert werden, so dass die Chips nicht nur kleiner, sondern auch leistungstärker wurden. Unglücklicherweise wird dieser Fortschritt in den nächsten Jahren an physikalische Grenzen stoßen. Neue Konzepte, wie zum Beispiel molekulare Transistoren, werden benötigt, um eine weitere Verbesserung unserer Computer zu erreichen, da sich die konventionellen Transistoren bald nicht mehr weiter verkleinern lassen.

Ein herkömmlicher Transistor funktioniert etwa wie ein Staudamm (Abb. 1): Ist er geschlossen, staut sich das Wasser auf der Seite A (Kollektor) und der Wasserstand auf der Seite B (Emitter) liegt auf einem niedrigeren Niveau. Erhält der Staudammwart (Basis) ein bestimmtes

¹ Für diesen Beitrag wurde A. Julia Stähler 2008 mit dem Klaus Tschira Preis für verständliche Wissenschaft ausgezeichnet (www.klaus-tschira-preis.info).

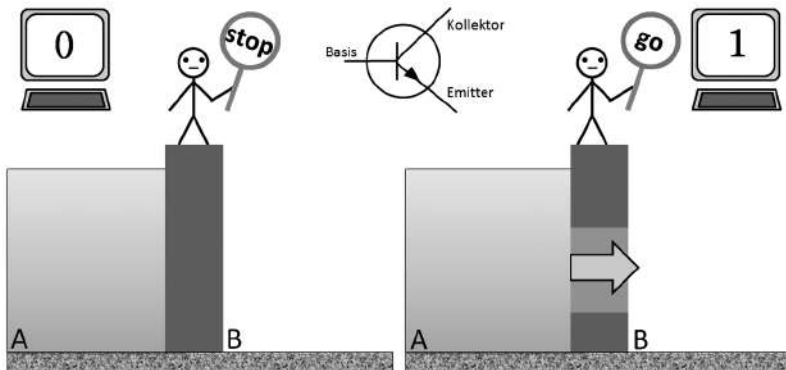


Abbildung 1: Funktionsweise eines herkömmlichen Transistors. Solange die Basis kein Signal erhält, fließt kein Strom. Erst wenn der Transistor geschaltet wird, fließen Elektronen zum Emitter. Dies ist vergleichbar mit einem Staudamm. Links: Ist er geschlossen, staut sich das Wasser auf der Seite A an und der Wasserstand auf der Seite B liegt auf einem niedrigeren Niveau. Rechts: Erhält der Staudammwart (Basis) ein Signal, öffnet er den Damm und das Wasser kann fließen.

Signal, öffnet er den Damm und das Wasser kann fließen. Ganz ähnlich sammeln sich im Transistor auf der Kollektorseite A Elektronen an, solange die Basis kein Signal erhält. Wird der Transistor geschaltet (geöffnet), können sie zur Emittenseite B fließen. Grundvoraussetzung für einen funktionstüchtigen Transistor ist es, dass die Seiten A und B durch eine undurchdringliche Barriere (Staudamm) getrennt sind, denn Informationen werden in einem Computer binär verarbeitet: Ein Transistor ist entweder auf oder zu, was vom PC als »1« oder »0« interpretiert wird. Im Fall eines undichten Staudamms wäre solch eine eindeutige Zuordnung nicht mehr möglich. Die stete Verkleinerung der Transistoren auf wenige millionstel Millimeter (Nanometer) führt dazu, dass die Barriere, die die Seiten A und B voneinander trennt, immer durchlässiger für Elektronen wird (Abb. 2). Quantenmechanische Effekte treten auf, die zu so genannten Leckströmen führen oder, in anderen Worten, der Staudamm wird undicht. Die durchschnittliche Zeit, die das Wasser auf der Seite A des Staudamms gehalten werden kann, ist eine charakteristische Größe für die Durchlässigkeit des Damms und abhängig von seiner Beschaffenheit (zum Beispiel Dicke). Im Fall der Elektronen bezeichnet man diese Zeit als ihre Lebens-

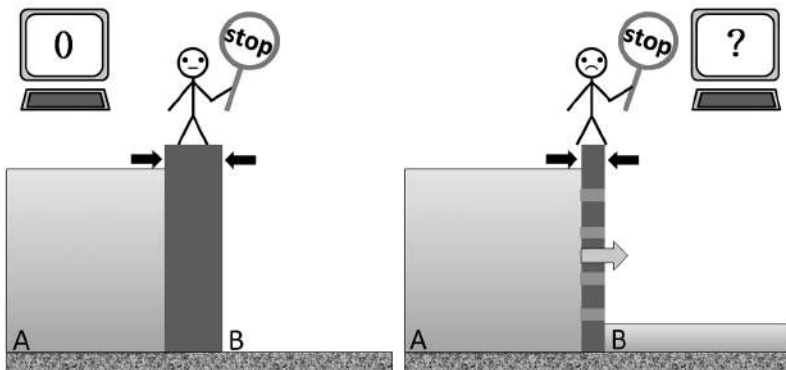


Abbildung 2: Physikalische Grenze der Verkleinerung von Transistoren. Links: Ein Transistor ist entweder auf oder zu, was vom PC als »1« oder »0« interpretiert wird. Rechts: Die Verkleinerung von Transistoren (schwarze Pfeile) führt dazu, dass die Barriere für Elektronen immer durchlässiger wird. Der Staudamm wird undicht, die eindeutige Zuordnung von »1« und »0« ist nicht mehr möglich.

dauer. Sie beschreibt, wie lange ein durchschnittliches Elektron von der Barriere davon abgehalten wird, von der Seite A zur Seite B zu wechseln.

In meiner Arbeit beschäftige ich mich eingehend mit solchen Elektronentransferprozessen, um über ein besseres Verständnis der zugrunde liegenden Prozesse an der Weiterentwicklung von molekularen Transistoren mitzuwirken. So konzentrierten mein Vorgänger Cornelius Gahl und ich uns unter anderem auf den Elektronentransfer an Eis-Metall-Grenzflächen, genauer gesagt, den Transfer von Elektronen durch extrem dünne Eisschichten (schmäler als zwei millionstel Millimeter) in ein Metall. Hierbei entspricht die Eisschicht dem oben erwähnten Staudamm und das Metall dem Fluss, der hinter dem Damm liegt. Die Messung der Lebensdauer der Elektronen in der Eisschicht (das heißt die mittlere Zeit, die sie vom Staudamm zurückgehalten werden) ist eine experimentelle Herausforderung, da diese nur winzige Bruchteile von Sekunden (Femtosekunden) beträgt. Ihre Bestimmung durch bloße Messung der Änderung der elektrischen Spannung zwischen Seite A (Eisschicht) und Seite B (Metall) ist nicht realisierbar. Glücklicherweise ist es heute mit Hilfe von ausgeklügelten Lasern

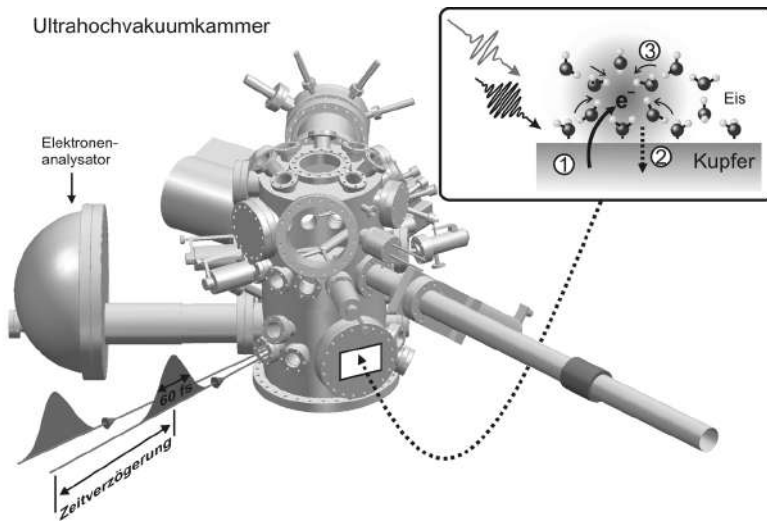


Abbildung 3: Experiment zur Untersuchung von Elektronentransferprozessen an Eis-Metall-Grenzflächen. Die Femtosekunden-Laserpulse werden in eine Ultrahochvakuumkammer geschickt, wo die Elektronen mit einem Analysator nachgewiesen werden. Der erste Laserpuls ist der Startschuss für das Experiment und der zweite »fotografiert« den Zustand, in dem sich die Elektronen befinden. Die Änderung der Zeitverzögerung zwischen dem Anrege- und dem Abfragepuls erlaubt das Mitverfolgen des Weges der Elektronen hin zum Gleichgewicht. Kasten: Elementare Prozesse an Eis-Metall-Grenzflächen nach Anregung mit einem Laserpuls: Elektronen (e^-) werden aus dem Metall in das Eis »geschossen« (1) und beginnen sofort, in das Metall zurückzutransferieren (2). Gleichzeitig orientieren sich die Wassermoleküle wegen der negativen Ladung des Elektrons um (3), was die Elektronen länger im Eis hält.

möglich, sehr kurze Lichtblitze zu erzeugen, deren Dauer vergleichbar ist mit der Lebensdauer von Elektronen an Grenzflächen. Mit einem solchen Laserpuls wird im Experiment das eisbedeckte Metall beleuchtet, was dazu führt, dass Elektronen in die Eisschicht »geschossen« werden (Prozess (1) in Abb. 3). Dieser erste Laserpuls ist in gewisser Weise der Startschuss für das eigentliche Experiment: Sobald Elektronen im Eis sind, beginnen auch schon die ersten, in das Metall zurückzutransferieren (Prozess (2) in Abb. 3), weil die Eisschicht keine unüberbrückbare Barriere für sie darstellt (wie der undichte Damm).

Wie Salz im Nudelwasser

Parallel zum Ladungstransfer findet noch ein weiterer Prozess statt, die so genannte Solvatisierung der Elektronen. Vom lateinischen Wort *solvere* (= lösen) abgeleitet, beschreibt der Begriff Solvatisierung das Auflösen der Elektronen im Eis, ganz ähnlich wie sich Salz im Nudelwasser auflöst. Dahinter steckt eine wichtige Eigenschaft des Wassermoleküls, welches an einem Ende eher positiv, am anderen eher negativ geladen ist. Das plötzliche Auftauchen der negativen Ladung des Elektrons im Eis führt dazu, dass sich die Wassermoleküle mit ihrem positiven Ende zum Elektron ausrichten (Prozess (3) in Abb. 3), weil sich gegensätzliche Ladungen anziehen. Durch diese Umorientierung kommt es zu einem erhöhten »Wohlfühleffekt« des Elektrons, denn zuvor waren noch etliche negative und damit abstoßende Enden der Moleküle auf es gerichtet. Diese Solvatisierung der Elektronen hat zur Folge, dass sie länger in der Eisschicht bleiben, da die Umorientierung der Moleküle die Barriere zwischen Elektronen und Metall verstärkt. Anders ausgedrückt, wird der Staudamm immer dichter, je länger der Stausee gefüllt ist. Der Fluss spült nach und nach Kies an den Staudamm, so dass die undichten Stellen mit der Zeit gestopft werden.

Beide Prozesse, der Elektronentransfer und die Solvatisierung, werden im Experiment mit Hilfe eines zweiten Laserpulses beobachtet. Dieser zweite Lichtblitz trifft um ein Zeitintervall versetzt nach dem ersten auf die Eis-Metall-Grenzfläche und »fotografiert« den *Status quo* der Elektronen, ihre Anzahl und den Grad ihrer Solvatisierung. So wie ein Videofilm aus einer Aneinanderreihung von Fotos besteht, lässt sich nun auch das gesamte »Leben« der Elektronen verfolgen: Ändert man das Zeitintervall zwischen erstem und zweitem Laserpuls, »filmt« man das Verhalten der Elektronen in Echtzeit. Auf diese Weise kann man zeigen, wie die Elektronen nach ihrem Einspeisen in die Eisschicht nach und nach durch die Wassermoleküle stabilisiert werden (Abb. 4 unten) und gleichzeitig kontinuierlich in das Metall zurückkehren (Abb. 4 oben), bis schließlich nach einer Pikosekunde (einer billionstel Sekunde) keine Elektronen mehr im Eis sind.

Solch kurze Lebensdauern sind natürlich nicht ausreichend, bedenkt man, dass die Barriere eines Transistors genügend dicht sein muss, um eine eindeutige Zuordnung seines Zustands (offen/geschlossen) zu ermöglichen. In meiner Doktorarbeit konnte ich zeigen, dass die Beschaffenheit der Molekülschicht, das heißt die Art, Anzahl und An-

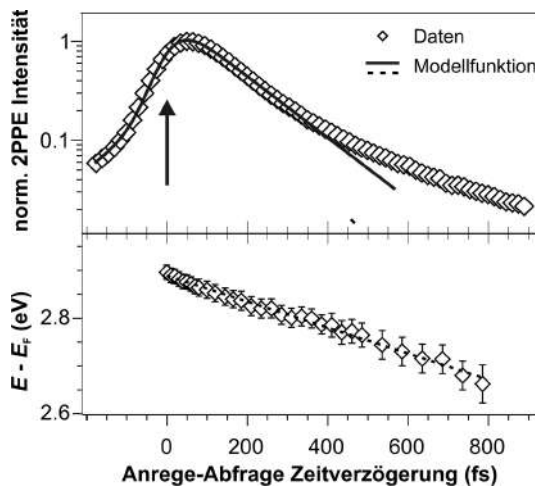


Abbildung 4: Ladungstransfer und Solvatisierung (Daten aus Gahl et al. 2002). Oben: Die Besetzung des solvatisierten Elektronenzustands nach der Anregung (Pfeil, Prozess (1) in Abb. 3) nimmt zunächst durch Ladungstransfer zurück in das Kupfer exponentiell ab (schwarze Kurve, Prozess (2) in Abb. 3). Nach etwa 300 fs beginnt die zusätzliche Abschirmung durch die Wassermoleküle zu wirken, und der Rückgang der Besetzung verlangsamt sich. Unten: Nach dem ersten Laserpuls ($t=0$ fs) sinkt das Energieniveau der Elektronen ($E-E_F$) in der Eisschicht durch die Umorientierung der Wassermoleküle ab (»Wohlfühleffekt«, Prozess (3) in Abb. 3).

ordnung der Moleküle, die zwischen dem Elektron und dem Metall liegen, großen Einfluss auf die Lebensdauer der Elektronen hat. Bemerkenswerterweise ändert sich das Verhalten der Elektronen besonders stark, wenn bestimmte Modifikationen der Anordnung der Wassermoleküle vorgenommen werden. So lässt sich die Lebensdauer der Elektronen durch Kristallisierung der Eisschicht, das heißt durch hohe Ordnung in der Struktur des Eises, von einigen hundert Femtosekunden auf mehrere Minuten (!) verlängern (Abb. 5 links). Das Einbringen dieser Elektronen in die Eisschicht und die darauf folgende Solvatisierung finden nach wie vor innerhalb weniger Femtosekunden statt, ist aber so viel effizienter, dass die Elektronen nun minutenlang in der Schicht überleben. In gewisser Weise werden die Löcher im Staudamm deutlich besser und schneller gestopft, so dass er dicht ist, bevor ein nennenswerter Anteil der Elektronen hindurchfließen kann

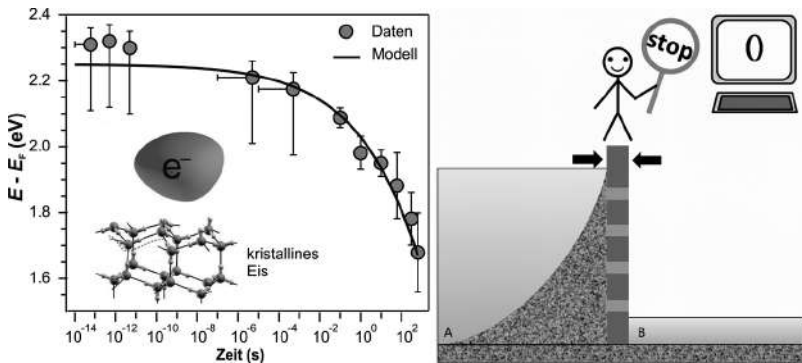


Abbildung 5: Verlängerung der Lebensdauer von Elektronen in der Eisschicht. Links: Bei hoch geordnetem, kristallinem Eis verlängert sich die Aufenthaltsdauer der Elektronen im Eis von Femtosekunden auf mehrere Minuten durch effizientere Abschirmung. So kann die Solvatisierung über 17 Größenordnungen der Zeit verfolgt werden. (Reprint with permission from Bovensiepen et al. Copyright 2009 American Chemical Society) Rechts: Die durch Solvatisierung verlängerte Lebensdauer entspricht dem Abdichten des undichten Staudamms durch angeschwemmtes Geröll.

(Abb. 5 rechts). Dieses Ergebnis ist äußerst bemerkenswert: Die Lebensdauer der Elektronen erhöht sich durch die Kristallisierung um ein 100 000 000 000 000-faches (100 Billionen), vergleichbar mit einer Eintagsfliege, die zwanzigmal länger lebt als das Universum.

Fazit

Zusammenfassend bleibt zu bemerken, dass selbst extrem dünne Molekülschichten dazu in der Lage sind, Elektronen äußerst effektiv von einer Metallgrenzfläche fernzuhalten, wenn ihre Struktur den Anforderungen entsprechend gewählt wurde. Diesen Vorteil gegenüber herkömmlichen Isolator- oder Halbleiterschichten in klassischen Transistoren verdanken sie dem Prozess der Solvatisierung, der dazu führt, dass sich der Staudamm nach und nach selbst verstärkt. Die Realisierung von Computerchips aus Eis ist allerdings höchst unwahrscheinlich, nicht zuletzt, weil die oben beschriebenen Experimente Temperaturen unter minus 200 Grad Celsius voraussetzen. Dessen

ungeachtet zeigen die beschriebenen Untersuchungen jedoch, dass eine weitere Verbesserung von Computerchips mittels molekularer Schaltelemente nicht nur prinzipiell möglich, sondern äußerst vielversprechend ist.

Literatur

- Bovensiepen U., Gahl C., Stähler J., Bockstedte M., Meyer M., Baletto F., Scandolo S., Zhu X.Y., Rubio A., Wolf M. (2009) A dynamic landscape from femtoseconds to minutes for excess electrons at ice-metal interfaces. *Journal of Physical Chemistry C* 113 (3), S. 979.
- Gahl C., Bovensiepen U., Frischkorn C., Wolf M. (2002) Ultrafast dynamics of electron localization and solvation in ice layers on Cu(111). *Physical Review Letters* 89, S. 107402.
- Stähler J., Bovensiepen U., Wolf M. (2010) Electron dynamics at polar molecule-metal interfaces: Competition between localization, solvation, and transfer. In: *Dynamics at Solid State Surfaces and Interfaces, Part I: Current Developments*, hrsg. von U. Bovensiepen, H. Petek, M. Wolf, Berlin (Wiley-VCH).
- Stähler J., Meyer M., Bovensiepen U., Wolf M. (2011) Solvation dynamics of surface-trapped electrons at NH_3 and D_2O crystallites adsorbed on metals: From femtosecond to minute timescales. *Chemical Science* 2 (5), S. 907.

Markus Hartl, Jan-W. Kellmann
und Ian T. Baldwin

Ökologische Gentechnik

Wie Gentechnik und Freilandversuche zu einem
besseren Verständnis von Ökosystemen beitragen¹

Das 20. Jahrhundert war geprägt von politischen Umwälzungen, katastrophalen und epochalen Ereignissen, von unglaublichem technischen Fortschritt und rasendem Wissenszuwachs. Es verwundert daher nicht, dass auch die Biologie in diesem Jahrhundert ihre eigene Revolution erlebt hat. Die Aufklärung der Zusammensetzung und Struktur von Desoxyribonukleinsäure (DNS oder DNA), der Trägerin der Erbinformation, und die Entschlüsselung des genetischen Codes bildeten den Grundstein für die moderne Molekularbiologie. Die Entwicklung besserer und vor allem schneller Sequenziertechnologien erlaubte es gegen Ende des 20. Jahrhunderts, die komplette Erbinformation von Organismen zu entziffern, was 2003 in der Veröffentlichung des menschlichen Genoms gipfelte.

In den letzten 30 Jahren haben sich die elektronischen Datenbanken mit DNA-Sequenzen der verschiedensten Lebensformen rasant gefüllt. Dieser enorme Wissenszuwachs und die phantastische Möglichkeit, zum ersten Mal im Buch des Lebens Buchstabe für Buchstabe lesen zu können, sind von so großer Bedeutung, dass in diesem Zusammenhang gerne von der »genomischen Revolution« gesprochen wird. Die Möglichkeit, den genetischen Code zu lesen, war ein Meilenstein für das Verständnis des Lebens, da die Biologie die Mittel bekam, zu verstehen, wie die einzelnen Gene als Erbinformation den Aufbau und die Funktion von Zellen und Lebewesen bestimmen, wie Leben also »funktioniert«.

Allerdings muss an dieser Stelle mit einem gängigen Missverständnis aufgeräumt werden. In den Medien wird gern von »Entschlüsselung« des Erbguts gesprochen, wenn Wissenschaftler die genetische Sequenz eines Gens oder, wie im Falle des Menschen, eines gesamten Genoms

¹ Vortrag beim XLAB Science Festival 2010.

herausfinden. Das erweckt oft den Eindruck, dass die Forscher bereits verstanden oder eben »entschlüsselt« hätten, welche Information in dieser Sequenz steckt. DNA zu sequenzieren, bedeutet jedoch nur, die Abfolge der vier Basen A, T, G und C in einem DNA-Molekül zu bestimmen. Molekularbiologen haben so den genetischen Code zwar buchstabiert, aber das heißt nicht, dass sie auch die Bedeutung eines Wortes, das sich daraus ergeben könnte, und damit die Funktion eines Gens kennen. Die einfache Aneinanderreihung der vier Basen in einem DNA-Molekül ist nämlich zunächst eher bedeutungslos. Zwar können erfahrene Molekularbiologen schon an der DNA-Sequenz verschiedene mögliche Eigenschaften des betreffenden Erbmolekül-Abschnitts erkennen, aber zumeist bleiben das Prognosen, die ähnlich wie die Wettervorhersage zutreffen können oder auch nicht. Das bedeutet: Wenn man erkennen will, wozu ein Gen gut ist, reicht es nicht aus, seine Sequenz zu kennen. Die Sequenz ist nur die Grundlage für die eigentliche und meist langwierige Arbeit der heutigen Biologen: nämlich herauszufinden, welche Funktion der sequenzierte DNA-Abschnitt, das darauf vorhandene Gen und schließlich das Protein, für welches das Gen ja die Bauanleitung liefert, eigentlich hat. Kurzum, die genomische Revolution hat uns eine riesige Menge an Daten beschert, die es nun zu verstehen gilt.

Natürlich waren die Biologen bis zur Sequenzierung ganzer Genome alles andere als untätig, sodass schon viele Prozesse und Funktionen einzelner Gene und Proteine bekannt waren. Bei einer geschätzten Gesamtzahl von etwa 22 500 Genen im menschlichen Genom bleibt aber natürlich noch sehr viel Arbeit übrig.² Warum ist es eigentlich so schwierig, die Funktion eines Gens oder eines Proteins herauszufinden? Um dies besser zu erklären, greifen wir auf einen Trick zurück, den Biologen gern anwenden, wenn eine Sache kompliziert wird: Wir wechseln zu einem einfachen Modell und versuchen, das Problem erst im Kleinen zu verstehen, bevor wir uns mit komplexen Lebensformen beschäftigen. Einer der Lieblingsmodellorganismen der Genetik ist seit jeher die Bäckerhefe (*Saccharomyces cerevisiae*). Auch wir werden uns diesem Lebewesen nun etwas genauer widmen.

- 2 Das kann man auch daran erkennen, dass sogar die Anzahl der Gene des Menschen nicht genau bekannt ist. Je nach Definition, Auslegung und Analyse der vorhandenen Daten belaufen sich momentane Schätzungen auf besagte 22 500, allerdings mit einem Spielraum von plus/minus 2000 Genen.

Von der Hefe lernen

Die Bäcker- oder Bierhefe und der Mensch haben schon einen langen gemeinsamen Weg hinter sich. Bereits vor mehr als 5000 Jahren wurde Hefe zur Herstellung alkoholischer Getränke und zum Lockern des Brotteiges verwendet. Viel später hat dann auch die Genetik die Hefe für sich entdeckt, weil sie ein paar unschätzbare Vorteile im Laboralltag bietet: Sie lässt sich leicht kultivieren, vermehrt sich schnell, ist einzellig, besitzt einen Zellkern und darin ein Genom mit überschaubarer Größe. Diese Eigenschaften erlauben es, relativ einfach Prozesse zu untersuchen, die grundlegend für alle Lebewesen mit Zellkern sind, wie die Replikation des Erbguts, den Zellzyklus, Signalwege in der Zelle und vieles mehr. Außerdem haben sich Genetiker eine Vielzahl gentechnischer Methoden einfallen lassen, mit denen gezielt Änderungen am Erbgut vorgenommen und ihre Wirkungen untersucht werden können. Das Modellsystem Hefe war und ist daher ein großer Erfolg, und der Funktion vieler menschlicher Gene und Proteine ist man auf die Spur gekommen, indem man verwandte Gene in der Hefe untersuchte.

Es ist also nicht verwunderlich, dass das Genom der Hefe das erste eukaryotische Genom (also das Genom eines Lebewesens mit Zellkern) war, das vollständig sequenziert werden sollte. 1996 war es soweit; und damals steckte man sich auch das ehrgeizige Ziel, in den folgenden zehn Jahren alle 6000 Gene der Hefe in ihrer Funktion zu beschreiben. 2007 musste man allerdings eingestehen, dass dies vielleicht doch etwas zu optimistisch war. Da waren nämlich noch etwa 1200 Gene übrig, von denen man so gut wie gar nichts wusste. Wie kann das sein, wo es doch weltweit mehr fleißige Hefe-Wissenschaftler als Hefe-Gene gibt? Dafür gibt es sicherlich mehrere Gründe, und manche davon sind ganz einfach, wie z.B. dass manche Gene noch gar nicht so lange bekannt sind, weil sie erst vor kurzem in der Genomsequenz entdeckt wurden. Einer der Hauptgründe liegt jedoch ganz woanders: Seit Jahrzehnten arbeiten Forscher mit Hefe im Labor, sie lassen sie in Glaskolben und auf Petrischalen wachsen und haben somit viele Details über die Genetik, Physiologie und Biochemie der Hefe erfahren. Mensch und Hefe haben also schon eine lange Zeit nebeneinander im Labor verbracht, aber dabei hat man ein sehr wichtiges Detail vergessen: Der ursprüngliche und natürliche Lebensraum von Hefe ist gar nicht das Labor! Vor lauter Konzentration auf die junge und moderne Genetik und die Erforschung der Stoffwechselvorgänge

hat man es versäumt, sich darum zu kümmern, wo und unter welchen Bedingungen Hefe in der Natur wächst und gedeiht.

Hefen finden sich eigentlich fast überall: Im Boden und im Wasser, auf Pflanzen oder Insekten oder sogar unter extremsten Bedingungen in der Tiefsee oder in polaren Böden. Natürlich handelt es sich dabei nicht um die Bäckerhefe, sondern ihre Verwandten, aber es zeigt doch, wo so ein einzelliger Hefepilz zu überleben in der Lage sein kann, wenn er nicht gerade im Labor im Reagenzglas in optimierten Nährlösungen schwimmt. Die Bäckerhefe selbst treibt sich wohl eher auf reifen Früchten, in Blütennektar und deshalb auch auf Insekten herum. Auch wenn wir das alles heute noch gar nicht so genau wissen, ist doch eines klar: Die Lebensbedingungen auf z.B. einer reifen Weintraube unterscheiden sich grundlegend von denen im Labor und es ist durchaus vorstellbar, dass einige, wenn nicht viele der noch nicht näher beschriebenen Gene mit diesem zweiten, geheimnisvollen und eigentlichen Leben der Hefe zu tun haben.

Unser Modell-System Hefe zeigt uns also, was notwendig ist, wenn wir die Funktionsweise eines Lebewesens mitsamt seinem Genom verstehen wollen: Wir müssen versuchen, jede Lebensform mit ihrer Erbinformation im Zusammenspiel mit ihrer natürlichen Umgebung zu begreifen. Es reicht bei weitem nicht aus, sich im Labor mit Zellbiologie, Genetik und Biochemie zu beschäftigen, sondern wir müssen hinaus in die Natur, um die Ökologie der Lebewesen zu erforschen. Das sagt einem ja eigentlich auch der gesunde Menschenverstand, sollte man meinen, aber zur Verteidigung aller Laborbiologen und ihrer unschätzbar wertvollen Arbeit sei eines gesagt: Viele der heutigen Standardmethoden der Molekularbiologie wurden über Jahrzehnte entwickelt und verfeinert, und das war an sich schon eine schwierige Aufgabe. Ohne einfache Modellsysteme wie das der Hefe wären die moderne Biologie und der damit verbundene – auch medizinische – Fortschritt gar nicht möglich gewesen. Außerdem war und ist die Beherrschung vieler Methoden nicht gerade trivial, sodass man bereits im Labor damit vollständig ausgelastet ist und nicht viel Zeit bleibt, zusätzlich noch Forschung im Freiland zu betreiben. Dennoch scheint es, dass die Entwicklung und Vereinfachung vieler Methoden in der Vergangenheit uns nun erlauben, den nächsten Schritt zu machen und Labor- mit Freilandforschung zu kombinieren. Wie das aussehen könnte und welche Methoden dafür von Bedeutung sind, soll uns im nächsten Kapitel beschäftigen.

Warum Pflanzen noch Blätter haben

Am Max-Planck-Institut für chemische Ökologie in Jena beschäftigt man sich in der Regel wenig mit Hefe- oder Humangenetik. Hier steht vor allem eine Frage im Mittelpunkt: Welche chemischen Signale sind für die Interaktion und das Überleben von Pflanzen, Tieren und Mikroorganismen in ihrer natürlichen Umwelt von Bedeutung? Um diese Frage zu beantworten, arbeiten eine Vielzahl von Ökologen, Physiologen, Entomologen, Chemikern, Biochemikern, Molekularbiologen, Neurobiologen, Genetikern und Verhaltensforschern im Institut zusammen. Das Hauptaugenmerk liegt dabei auf der Interaktion von Pflanzen und Insekten, den beiden häufigsten und vielfältigsten makroskopischen Lebensformen auf unserem Planeten. Pflanzen sind außerdem insofern sehr interessante Lebewesen, weil sie, verglichen mit Tieren, einen Nachteil haben: Sie können vor Problemen und Gefahren nicht davonlaufen! Eine Pflanze ist in der Regel an demjenigen Ort fest verankert, an dem sie aus dem Samen gekeimt ist. Von diesem Zeitpunkt an ist die Pflanze einer Vielzahl von Umwelteinflüssen ausgesetzt, mit denen sie irgendwie fertig werden muss. Dabei unterscheidet man in der Regel zwei Gruppen von Umwelteinflüssen: 1. abiotische Faktoren, also alle Einflüsse der unbelebten Umwelt, wie z.B. Temperatur, Niederschlag, Nährstoffverfügbarkeit, Sonneneinstrahlung etc., und 2. biotische Faktoren, also alle Einflüsse und Interaktionen von und mit anderen Lebewesen (Abb. 1). Diese können sowohl negativ für die Pflanze, wie z.B. Fraß durch Insekten oder Säugetiere oder Infektion durch pathogene Viren und Mikroorganismen, als auch positiv sein, wie z.B. Bestäubung der Blüten durch Insekten, Verfügbarmachung von Nährstoffen durch symbiotische Mikroorganismen und vieles mehr. Auch wenn wir uns am Max-Planck-Institut für chemische Ökologie vor allem für die biotischen Faktoren interessieren, so wissen wir, dass das eine künstliche Trennung ist und sich in der Wirklichkeit sämtliche Faktoren gegenseitig beeinflussen können. Trockenheit verändert z.B. auch die Bedingungen für Mikroorganismen im Boden, oder hohe Feuchtigkeit befördert die Verbreitung von Pilzen. Dennoch hilft uns diese Unterscheidung, einen Überblick über die vielen Umweltfaktoren zu bekommen, und dient dazu, die Prozesse einzeln kennen zu lernen, bevor man versucht, alles auf einmal zu untersuchen.

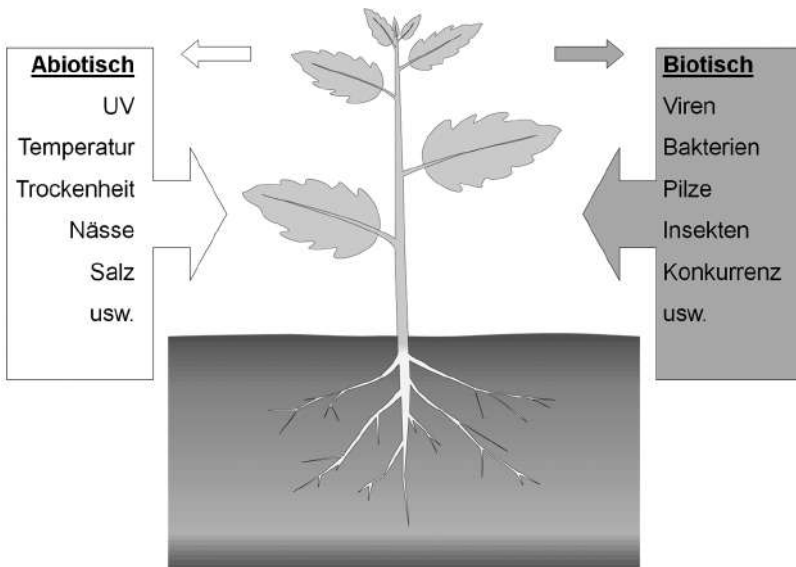


Abbildung 1: Beispiele abiotischer und biotischer Umweltfaktoren, die für das Leben und Überleben von Pflanzen von Bedeutung sind.

Wir wollen uns aber nun auf die Interaktionen von Pflanzen mit ihrer belebten Umwelt konzentrieren, und zwar auf eine eher unangenehme Angelegenheit: das Gefressenwerden. Wie bereits erwähnt, kann eine Pflanze vor ihren Fraßfeinden (Herbivoren) nicht davonlaufen. Welche Möglichkeiten hat sie, trotzdem zu überleben? Oder anders gefragt: Warum haben die Pflanzen um uns herum noch Blätter und sind nicht völlig entlaubt? Die Evolution hat Pflanzen im Laufe ihrer etwa 400 Mio. Jahre langen gemeinsamen Entwicklung mit ihren Fraßfeinden mit einer Vielzahl von Verteidigungsstrategien ausgestattet. Dazu zählen z.B. mechanische Abwehrmechanismen wie Dornen, direkte chemische Verteidigung mit Giftstoffen, aber auch indirekte Verteidigungsstrategien durch Anlockung der Feinde der Fraßfeinde – doch dazu kommen wir im Detail später. Jede Entwicklung eines Abwehrmechanismus hat einen entsprechenden Selektionsdruck auf die Fraßfeinde ausgeübt, sodass diese Gegenstrategien entwickelt haben, um z.B. die Gifte unschädlich zu machen, zu vermeiden oder sogar für eigene Zwecke zu verwenden. Ein derartiger evolutiver Prozess gegenseitiger Anpassung wird deshalb auch Co-Evolution genannt und

könnte einer der Gründe dafür sein, warum wir heute eine so große Vielfalt von Pflanzen und Insekten auf der Erde vorfinden.³ Der wiederkehrende Zyklus von Anpassung und Gegen-Anpassung führte zu immer neuen Eigenschaften und damit letztlich verschiedenen Arten und ist natürlich nicht nur auf Pflanzen-Herbivoren-Interaktionen begrenzt, sondern betrifft jegliche Form von ökologischer Interaktion. Die Interaktionen zwischen Pflanzen und Insekten sind dabei für Biologen allerdings besonders spannend, weil Insekten die artenreichste Tiergruppe der Erde darstellen und sich über Millionen von Jahren eine unüberschaubare Anzahl von faszinierenden, co-evolutionären Beziehungen zwischen Insekten und Pflanzen entwickelt hat.

Das Ergebnis und Protokoll der Evolutionsgeschichte eines jeden Lebewesens ist letztlich sein Erbgut, sein Genom, womit wir wieder beim ersten Kapitel wären. Wenn wir also verstehen wollen, wozu alle die Gene in einer Pflanze oder in einem Insekt gut sind, dann müssen wir immer die natürliche Umwelt des Lebewesens mit einbeziehen – eben wie bei der Hefe. Nun kommt allerdings das Problem: Statt mit etwa 6000 Genen wie in der Hefe haben wir es in Pflanzen mit weitaus mehr zu tun. Das Lieblingsobjekt der Pflanzengenetiker ist ein eher unscheinbares Pflänzchen mit dem Namen *Arabidopsis thaliana*, zu Deutsch: Acker-Schmalwand. Sie hat zwar ein eher kleines Genom, das im Jahr 2002 sequenziert wurde, aber es enthält nach momentanem Wissen etwa 27 400 Gene – also ungefähr 5000 mehr als beim Menschen. Natürlich ist es bei *Arabidopsis* so ähnlich wie bei der Hefe. Die Wissenschaft war schon sehr fleißig und man weiß heute sehr viel darüber, wie diese Pflanze wächst und sich vermehrt und welche Gene dabei relevant sind. Dennoch tapen wir bei einer Vielzahl von Genen noch im Dunkeln und die Vermutung liegt nahe, dass wir ähnlich wie bei der Hefe noch nicht die richtigen Umweltbedingungen gefunden haben, unter denen sich die Funktion dieser Gene offenbart. An dieser Stelle müssen wir nun endlich eine Sache näher beleuchten: Wie untersucht man eigentlich die Funktion eines Gens?

- 3 Wir Menschen haben übrigens von dieser Vielfalt bisher sehr profitiert. Nicht nur, dass sie die Grundlage für unsere reiche Speisekarte und viele Industrierohstoffe ist, sondern die Vielzahl an chemischen Strukturen, die sich in Pflanzen und anderen Lebewesen finden lässt, ist bis heute auch eine wichtige Basis und ein Hoffnungsträger für lebenswichtige Pharmazeutika.

Wozu man Gentechnik verwenden kann

Seit den 1970er Jahren haben Molekularbiologen nicht nur, was die Sequenzierung von DNA, sondern auch, was ihre gezielte Veränderung angeht, unglaubliche technologische und methodische Fortschritte erzielt. Das bedeutet, dass es heute in der Regel kein Problem mehr ist, von einer bekannten Gensequenz eine Kopie zu erstellen, sie zu verändern oder mit anderen Sequenzen zu kombinieren und sie in anderen Organismen funktionstüchtig wieder einzusetzen. Die dazugehörigen gentechnischen Methoden sind seit über zwei Jahrzehnten Standardmethoden der Biologie und Medizin und aus dem Laboralltag nicht mehr wegzudenken. So werden Gene z.B. isoliert und in Bakterien eingesetzt. Die Bakterien produzieren mit dieser Information ein Protein und der Forscher kann damit weitere Versuche machen, um die Eigenschaften des Proteins näher zu untersuchen.⁴ Proteine können auch mit Farbstoffen ausgestattet werden, damit man erkennen kann, wann in welchen Zellregionen oder in welchen Geweben sie ihre Funktion erfüllen, oder so verändert werden, dass man erfährt, ob und wann ein Protein mit einem anderen interagiert. Das sind nur ein paar Beispiele der gentechnischen Werkzeugkiste und man sollte vielleicht auch erwähnen, dass das alles viel weniger mit Frankenstein und Spektakel zu tun hat, als es uns *Jurassic Park* und ähnliche Geschichten weismachen wollen. Viel eher handelt es sich um präzises, methodisch ausgefeiltes, teilweise langwieriges und durch Behörden überwacht Handwerk, von dem, solange man nicht an gefährlichen Organismen oder deren Kombinationen arbeitet, keine Gefahr ausgeht. Die Organismen, die zum Einsatz kommen, sind in der Regel harmlose Bakterien, Hefen oder Zelllinien, die außerhalb des Labors überhaupt nicht lebensfähig sind. Wie so oft liegen das Risikopotenzial und die Gefahren also nicht direkt in der Technik, sondern bei den Menschen, die sie verwenden, und in den Absichten, die sie dabei hegen.

Eine bestimmte DNA-Sequenz in das Genom eines Bakteriums zu transferieren, ist relativ einfach. Dasselbe in Pflanzen oder anderen

4 Auf diese Weise werden übrigens auch Medikamente hergestellt. Seit bald 30 Jahren wird z.B. menschliches Insulin aus gentechnisch veränderten Bakterien gewonnen, das in der Herstellung billiger und für Diabetes-Patienten besser verträglich ist als das davor eher mühsam gewonnene Schweine- und Rinderinsulin.

mehrzelligen Organismen zu machen, ist hingegen schon etwas schwieriger. Glücklicherweise entdeckte man, dass das im Erdboden vorkommende Bakterium *Agrobacterium tumefaciens* genau das kann. Es überträgt natürlicherweise bestimmte Gene in die Pflanze, um sie für sich »umzuprogrammieren« und Nährstoffe herstellen zu lassen. Diese Eigenschaft haben sich die Pflanzengenetiker zu Nutze gemacht und verwenden das Bakterium nun dazu, anstelle der bakteriellen Gene die gewünschten Zielgene zu übertragen (z.B. um den Einfluss eines Gens auf die Pflanze zu untersuchen). Das Bakterium hat auch die Eigenschaft, bei der Integration eines fremden Gens (Transgens) in das pflanzliche Genom Pflanzengene zufällig zu unterbrechen und sie damit funktionsuntüchtig zu machen. Der Pflanze fehlt dann sozusagen ein Gen, und man kann untersuchen, wie sich das Fehlen der Information auf die Pflanze auswirkt. Dieser »Unterbrechungseffekt« war bisher sehr hilfreich, weil alle Unterschiede, die im Vergleich mit einer unveränderten Pflanze auftreten, auf das Fehlen dieses einen unterbrochenen Gens zurückzuführen sein müssen. Der Effekt hat aber zwei entscheidende Nachteile: 1. Der Prozess ist zufällig. Man kann *Agrobacterium* nicht vorschreiben, welches Gen es unterbrechen soll. 2. Unterbricht man Gene, die notwendig für das Überleben der Pflanze sind, dann ist es unmöglich, eine lebende Pflanze zu erhalten, die man untersuchen könnte. Alles, was man dann weiß, ist, dass das Gen überlebensnotwendig ist, nicht aber, warum.

Glücklicherweise haben eine Reihe von Wissenschaftlern eine elegante Methode entwickelt, die dieses Dilemma umgeht und deren Entdeckung und Erforschung 2006 mit dem Nobelpreis für Medizin bedacht wurde (die Preisträger waren Andrew Fire und Craig Mellow). Es handelt sich dabei um die sogenannte RNA-Interferenz, kurz RNAi. RNAi ist ein natürlich vorkommender Prozess, mit dem die Zelle mRNA spezifisch abbauen und damit regulieren kann. Zur Erinnerung: Die mRNA (*messenger RNA* oder Boten-RNA) ist die von der DNA abgelesene Erbinformation eines bestimmten Gens, sozusagen die Bauanleitung, die dann in der Regel in ein Protein übersetzt wird (Abb. 2 a). Wird die mRNA zerstört, kann sie in kein Protein mehr übersetzt werden, das Gen wird damit in seiner Funktion ausgeschaltet oder stillgelegt (im Englischen spricht man von *Gene Silencing*; Abb. 2 b). Interessanterweise braucht man dazu nichts anderes als ein kurzes Stück eines doppelsträngigen RNA-Moleküls (= dsRNA, mRNA ist einzelsträngig), das in seiner Sequenz mit der des auszu-

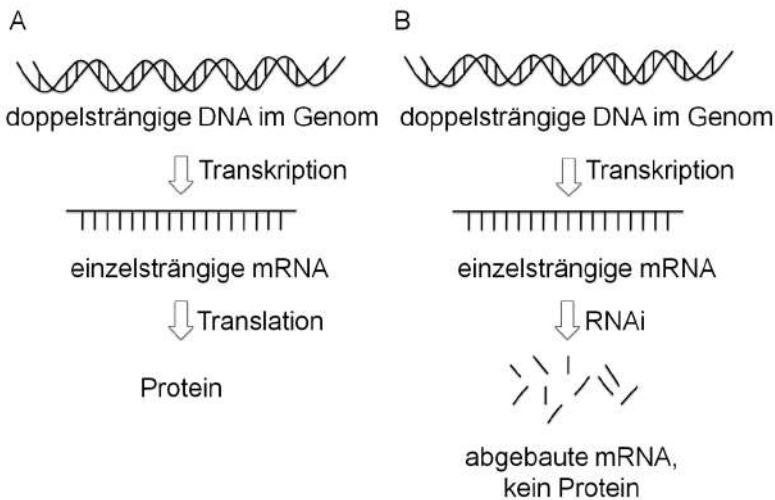


Abbildung 2: Genexpression und der Wirkmechanismus der RNA-Interferenz. A: der reguläre Ablauf von Transkription und Translation übersetzt die genetische Information eines Gens über die mRNA in ein spezifisches Protein. B: RNA-Interferenz zerstört gezielt und selektiv die mRNA und unterbindet dadurch die Produktion bestimmter Proteine und damit die Ausprägung des Gens.

schaltenden Gens identisch ist. Dieses Stück dsRNA dient einem bestimmten Enzym als Such-Muster, um die Ziel-mRNA aufzuspüren und abzubauen. Genau dies ist der Mechanismus, den man sich zu Nutze macht. Erzeugt oder bringt man in eine Zelle ein derartiges dsRNA-Molekül mit der Sequenz des gewünschten Gens ein, so aktiviert man damit den natürlichen RNAi-Ausschaltmechanismus, die mRNA wird abgebaut und das dazugehörige Gen inaktiviert. Es gibt verschiedene Methoden, wie man die dsRNA in einen Organismus hineinbekommt, aber wir wollen hier nicht zu sehr ins Detail gehen. Bei Pflanzen bedient man sich in der Regel der Möglichkeit, mit besagtem *Agrobacterium* eine Reihe von DNA-Sequenzen in das Genom der Pflanze zu transferieren, die dazu führen, dass die Pflanze die gewünschte dsRNA selbst produziert und somit das Gen permanent ausschaltet. Diese Methode vereint die beiden Vorteile, dass sie zielgerichtet das gewünschte Gen inaktiviert und dass sie auch bei Genen erfolgreich sein kann, die essenziell für das Überleben sind.

Nun haben wir also eine Werkzeugkiste beisammen, mit der wir Genfunktionen überprüfen können. Aber wo fängt man an und wie hat man sich das praktisch vorzustellen? Nehmen wir z.B. an, wir interessieren uns für Gen XY in *Arabidopsis*. Erste Experimente haben uns gezeigt, dass das Gen XY immer dann aktiviert wird, wenn die Pflanze von einem bestimmten pathogenen Pilz befallen wird. Ist das nun Zufall, eine Nebenwirkung der Infektion oder hat es etwas damit zu tun, dass die Pflanze versucht, auf die Infektion zu reagieren, sich also zu wehren? Um dieser Frage auf den Grund zu gehen, könnte man nun mittels RNAi Pflanzen herstellen, in denen das Gen XY unterdrückt ist. Infiziert man nun diese Pflanze mit dem Pilz und vergleicht sie mit einer ebenso infizierten, aber ansonsten unveränderten Pflanze (dem so genannten Wildtyp), könnte man sehen, ob sich diese Pflanze im Vergleich mit dem Wildtyp schlechter zur Wehr setzen kann, sich der Pilz schneller ausbreitet etc. Natürlich würde es an dieser Stelle erst spannend werden, man würde weiter ins Detail gehen und versuchen herauszufinden, auf welche Weise die Interaktion beeinflusst wird. Wichtig bleibt dabei allerdings stets der Vergleich der RNAi-Pflanze mit dem Wildtyp, weil alle beobachteten Unterschiede immer nur auf das Ausschalten des einen Gens zurückgeführt werden können.

Das war nun ein relativ einfaches Beispiel, das in der Realität, selbst wenn alles nach Plan läuft, dennoch zur Durchführung einige Jahre in Anspruch nehmen und z.B. einen Doktoranden für seine gesamte Doktorarbeit beschäftigen kann. Es war aber auch deswegen noch simpel, weil wir zu Beginn bereits wussten, dass das Gen XY wahrscheinlich etwas mit der Interaktion mit einem bestimmten Pilz zu tun hat. Genau das ist aber die Information, die oft so schwierig zu beschaffen ist. Stellen wir uns vor, der Vorversuch wurde falsch interpretiert und das Gen XY hat am Ende gar nichts mit der Pilzinfektion zu tun. Da kann der arme Doktorand noch so oft seine Pflanzen mit dem Pilz infizieren, er wird keine Unterschiede finden, weil das Gen nämlich – nehmen wir einmal an – etwas mit Toleranz gegenüber Kälte zu tun hat. Da der Doktorand seine Pflanzen aber immer schön unter Standardbedingungen bei gleicher, eher warmer Temperatur im Gewächshaus angezogen hat, wird er das nie herausfinden. Und hier befinden wir uns nun in demselben Dilemma wie die Hefegenetiker. Wir haben alle Möglichkeiten, die Gene einer Pflanze zu untersuchen, aber im Labor bzw. im Gewächshaus herrschen schlichtweg die falschen Bedingungen, um die Funktion vieler dieser Gene zu offenbaren. Und in der

Tat: Viele der transgenen Linien, bei denen mittels *Agrobacterium* ein Gen komplett unterbrochen wurde, zeigen unter Standard-Laborbedingungen keinen Phänotyp, also keine erkennbare Änderung zum Wildtyp. Nun könnte man natürlich versuchen, jeden einzelnen Faktor im Labor nachzubilden, aber bei der Vielzahl der möglichen Faktoren und ihrer Kombinationen (Abb. 1) würde das nahezu ewig dauern oder schlicht unmöglich sein, weil man entweder nicht alle Faktoren kennt oder es einfach zu viele sind. Wir beobachten hier sozusagen genau das umgekehrte Problem, mit dem sich Ökologen seit jeher herumschlagen: Ein Ökologe versucht nämlich, möglichst alle Faktoren in der Umwelt zu messen und zu erfassen, um ein bestimmtes beobachtetes Phänomen zu erklären, aber weil es so viele sind, ist es sehr schwierig, die entscheidenden Faktoren von den unwichtigen zu trennen. Man könnte also sagen: das Labor ist zu simpel, die Natur zu komplex.

Um diesem Problem zu begegnen, ist die Arbeitsgruppe Molekulare Ökologie am MPI für chemische Ökologie dabei, die Stärken beider Ansätze zu vereinen. Wie wäre es, wenn wir die Methode des Ausschaltens von Genen und den Vergleich mit unveränderten Pflanzen mit klassischen ökologischen Freilandexperimenten kombinierten? Es wäre dann nicht mehr notwendig, alle Umweltfaktoren zu kennen und zu simulieren, sondern die Natur selbst würde sich darum kümmern. Die Natur wäre dann sozusagen der Experimentator, der die Pflanzen mit allen Faktoren konfrontiert. Dem Wissenschaftler bleibt dann die Aufgabe zu beobachten, wie sich die Pflanze entwickelt, ob sie Schaden nimmt oder vielleicht besser wächst, mehr Fraßfeinde anlockt oder sie eher vertreibt, usw. Das erfordert eine gute Beobachtungsgabe, etwas Geduld und eine gewisse Flexibilität, aber es bietet den unglaublichen Vorteil, dass einem die Natur selbst zeigt, worin der Unterschied zwischen den Versuchs- und Vergleichspflanzen besteht. Diese Idee bildet das grundlegende Konzept für eine Reihe von Versuchen, die unsere Arbeitsgruppe an Pflanzen im Freiland vorgenommen hat, um die ökologische Bedeutung ausgesuchter Gene näher zu untersuchen. Wir möchten in der Folge zwei Fallbeispiele davon vorstellen, die einen Eindruck geben sollen, welche Vorteile diese Methode bietet. Dazu ist allerdings noch vorzuschicken, dass diese Versuche weder mit *Arabidopsis* noch mit bekannten Nutzpflanzen durchgeführt wurden, sondern mit auf den ersten Blick eher exotisch anmutenden Spezies: dem Wilden Tabak (*Nicotiana attenuata*) und dem Schwarzen

Nachtschatten (*Solanum nigrum*). Hinter dieser Auswahl steht allerdings keine merkwürdige Laune der Wissenschaftler, sondern genaues Kalkül. Schließlich geht es uns ja darum, die natürliche Situation zu verstehen: Welche Gene haben in einer natürlichen, wilden Pflanzenart welche Funktion? Alle unsere Nutzpflanzen nämlich, sei es Weizen oder Kartoffel, Mais oder Tomaten, wurden vom Menschen über Jahrtausende gezüchtet, um mehr Ertrag zu liefern, größere Früchte zu produzieren, besser zu schmecken usw. Bei diesem künstlichen, durch den Menschen gesteuerten Auswahlprozess sind aber – und dies bestätigen genetische Untersuchungen immer wieder – natürliche Eigenschaften verloren gegangen, die für die ursprüngliche Pflanze, als diese noch in der Wildnis allein bestehen musste, von Bedeutung waren. Nicht ohne Grund sind viele unserer heutigen Nutzpflanzensorten ohne Hilfe des Menschen durch den Einsatz von Spritzmitteln oder andere Maßnahmen nur bedingt überlebensfähig. Jeder, der schon einmal gesehen hat, was Schnecken mit einem Salatbeet anrichten können, weiß, wovon die Rede ist. Die Nutzpflanzen wurden sozusagen domestiziert und, man erlaube den Vergleich, es würde auch niemand ernsthaft auf die Idee kommen, das Leben von Wölfen in freier Wildbahn am Beispiel eines Zwergpudels zu untersuchen. Aus diesem Grund haben wir uns also für Wildpflanzen entschieden, und warum gerade für diese beiden bestimmten Arten, ist ebenso leicht erklärt. Beide sind nahe Verwandte von molekularbiologisch gut untersuchten Arten, nämlich einerseits dem Tabak (*Nicotiana tabacum*) und andererseits der Kartoffel (*Solanum tuberosum*) und der Tomate (*Solanum lycopersicum*). Das hat den unschätzbaren Vorteil, dass sich viele Methoden und Gen-Datenbanken, die für diese Nutzpflanzen erstellt wurden, relativ leicht auf die Wildarten übertragen lassen. Nun könnte man natürlich noch einwenden, dass es am sinnvollsten wäre, *Arabidopsis thaliana* zu verwenden, da doch über diese Pflanze bereits am meisten bekannt ist. Das ist zwar im Prinzip richtig, nur eignet sich *Arabidopsis* nur bedingt für die Untersuchung von Pflanzen-Herbivoren-Interaktionen. In der Natur wird *Arabidopsis* nämlich eher selten von Herbivoren heimgesucht, sodass die Pflanze nicht gerade eine optimale Wahl für unsere Untersuchungen darstellen würde. Im Gegensatz dazu werden unsere beiden ausgewählten Arten von einer Vielzahl verschiedener Herbivoren befallen, von denen manche wahre Spezialisten sind, weshalb sich, wie wir nun zeigen wollen, die Untersuchung ihrer co-evolutiven Anpassungen besonders lohnt.

Fallbeispiel 1: Wilder Tabak in Utah, USA

Lebensraum und Überlebensstrategien



Abbildung 3: Wildtyp (nicht gentechnisch veränderte Pflanze) des Wilden Tabaks (*Nicotiana attenuata*) an seinem natürlichen Standort im *Great Basin Desert* in Utah, USA (Foto: Celia Diezel). (Farbabbildung 2, S. 144)

Der Wilde Tabak, auch bezeichnet als *coyote tobacco*, ist eine im Westen Nordamerikas heimische Pflanze (Abb. 3). Den einstmals dort lebenden Völkern wie den Hopi oder Apachen diente die Pflanze als bevorzugter Rauchtobak – sicherlich aufgrund ihres hohen Gehalts an Nikotin, einem Gift, das die Pflanze bei Befall mit Pflanzenfressern mobilisiert. Mithilfe des wirkungsvollen Alkaloids wehrt der Tabak eine ganze Reihe von Schädlingen ab, und auch Kleinsäugern wird mindestens der Appetit an den nikotinhalten Blättern verdorben. Zusätzlich dient Nikotin als Bitterstoff im süßen Blütennektar. In wohldosierter Menge sorgt das Gift für kurze Aufenthalte und damit hohe Besucherfrequenzen durch pollen-

übertragende Kolibris und Schwärmermotten und optimiert so die Produktion neuer fruchtbarer Tabaksamen [1].

Um seine Art zu erhalten, hat der Wilde Tabak im Laufe der Evolution unterschiedliche Verteidigungslinien gegen seine Feinde aufgebaut. Besonders gut untersucht ist eine natürlich vorkommende Tabakpopulation im *Great Basin Desert* in der Nähe der Stadt St. George im Staat Utah (USA) nahe der Grenze zu Arizona. Seit 1996 betreibt das Max-Planck-Institut für chemische Ökologie dort eine Freilandstation auf dem Gelände des Naturschutzgebiets *Lytle Ranch*. Die große Becken-Wüste zeichnet sich durch steppenartige Gegenden aus mit einer je nach Region charakteristischen Flora und Fauna. Im Erdboden des *Lytle Ranch*-Gebiets befinden sich unzählige Samen des Wilden Tabaks, die immer dann konzentriert aufkeimen, sobald sie ein bestimmtes



Abbildung 4: Larve des Tabakswärmers (*Manduca sexta*) auf einem Blatt des Wilden Tabaks (Foto: Danny Kessler). (Farbabbildung 4, S. 146)

chemisches Signal aus der Umwelt erhalten. Dieses spezifische Signal entsteht nach Buschfeuern und ist im Rauch des verkohlten Holzes enthalten [2]. Innerhalb weniger Tage keimen zahlreiche Tabakpflanzen und bedecken mit ihren Rosettenblättern den Boden fast vollständig – ein gefundenes Fressen für herbivore Insektenarten, besonders für diejenigen, denen das Nikotingift nichts mehr anhaben kann. Zu diesen Pflanzenfressern gehört die Larve des Tabakswärmers *Manduca sexta* (Abb. 4). Die Spezies ist gegen Nikotin weitgehend resistent, und es bedarf nur weniger Raupen, um eine ausgewachsene Tabakpflanze innerhalb kurzer Zeit vollständig zu entlauben. Dennoch hat sich *Nicotiana attenuata* trotz derartiger gefährlicher Feinde am Standort erhalten; wahrscheinlich seit schon vielen Jahrhunderten (über-)lebt die Pflanzenart aus der Familie der Nachtschattengewächse in Utah im ökologischen Gleichgewicht mit ihren Fraßfeinden, auch mit denen, die sie nicht (mehr) direkt vergiften kann.

Wie also schafft es der Wilde Tabak, dass die Population eines Schädlings wie *Manduca*, der die direkte Verteidigungslinie einer Nikotinvergiftung überwinden kann, nicht überhandnimmt und die Spezies an ihrem Standort ausrottet? Bei der *Nicotiana-Manduca*-Interaktion handelt es sich um eine der ersten Pflanze-Insekt-Wechsel-

wirkungen, bei der die so genannte indirekte Verteidigung von Pflanzen gegen Schädlinge entdeckt und genau studiert wurde, sowohl im Freiland als auch in den Gewächshäusern des Max-Planck-Instituts für chemische Ökologie gleich nach dessen Gründung im Jahr 1996. Indirekte Verteidigung heißt, dass ein Schädling von einer Pflanze abgewehrt wird, indem sie, ökologisch ausgedrückt, eine dritte trophische Ebene rekrutiert: Sie lockt die Feinde ihrer Fraßfeinde herbei, und dies durch die Abgabe von flüchtigen Duftstoffen aus ihren grünen Blättern; beim Wilden Tabak handelt es sich um einen komplizierten Duftstoffcocktail [3]. Schlupfwespen oder Raubwanzen, je nach Standort der Pflanzenpopulation, werden durch solche Duftstoffe angezogen. Die Wespenweibchen legen ihre befruchteten Eier in die lebendigen *Manduca*-Raupen ab, und die kurz danach schlüpfenden Larven parasitieren den Pflanzenfresser. Raubwanzen der Gattung *Geocoris* wiederum interessieren sich besonders für frisch geschlüpfte Raupen und die Eier von *Manduca sexta*, die auf den Blattunterseiten der Tabakblätter kleben. Die Wanzen werden von einem speziellen Duftstoff, nämlich (E)-2-Hexenal, angelockt, den der Tabak am Ende gar nicht selbst bereitstellt, sondern der aus einem pflanzlichen Vorläufermolekül, nämlich (Z)-3-Hexenal, gebildet wird. Die Umwandlung zu (E)-2-Hexenal wird durch ein Enzym im Raupenregurgitat – das ist eine Mischung aus Speichel und Darminhalt – vorgenommen. Die junge Raupe verrät sich selbst also ungewollt an ihren Feind, indem sie einen Duftstoff produzieren hilft, auf den die Wanzen im wahrsten Sinne des Wortes »fliegen« [4].

Alle diese Vorgänge zeigen in eindrucksvoller Weise, dass mithilfe von giftigen Molekülen, wie beispielsweise dem Nikotin, aber auch mit leicht flüchtigen Molekülen, die von den Blättern des Tabaks abgegeben werden, eine erfolgreiche und vor allem effektive und ganz natürliche Schädlingsbekämpfung seitens der Pflanze vollzogen werden kann. Wie aber »weiß« der Tabak eigentlich, dass er von einer Raupe befallen ist? Wie funktioniert die Signalgebung innerhalb der Pflanze, also vom Raupenbiss in das Blatt bis hin zur Genexpression?

Steuerung der molekularen Abwehrmechanismen: Lipoxygenasen und Jasmonat

Über die molekularen Signalwege im Organismus Pflanze bei Befall durch Schädlinge oder mikrobielle Pathogene ist in den letzten dreißig

Jahren ein enormer Erkenntnisgewinn erzielt worden. Besonders seit Beginn der 1990er Jahre konnten dank dem Einsatz transgener Versuchspflanzen die einzelnen Schritte zellulärer Signalketten von der Plasmamembran bis in den Zellkern nachvollzogen werden. Regulatorische Bereiche (Promotoren) derjenigen Gene, die in die verschiedenen pflanzlichen Verteidigungsmechanismen involviert sind, konnten sehr genau durch den Einsatz chimärer Reportergenkonstrukte charakterisiert werden. Reportergene wie beispielsweise das *Green Fluorescent Protein*, das inzwischen in vielen verschiedenen Farbvariationen zur Verfügung steht, ermöglichten die Lokalisierung von verteidigungsspezifischen Proteinen innerhalb der Zelle oder in unterschiedlichen pflanzlichen Geweben. Dank einer anderen Gruppe chimärer Genkonstrukte, die entweder die Überexpression oder aber das Abschalten eines bestimmten Abwehr-Enzyms ermöglichen, konnte in Experimenten mit transgenen Versuchspflanzen dessen Rolle innerhalb eines gegebenen Verteidigungsmechanismus exakt bestimmt werden. In der Mehrzahl der Experimente mit transgenen Pflanzen des Wilden Tabaks wurden verteidigungsrelevante Gene abgeschaltet oder deren Expression deutlich unterdrückt. Das Stilllegen der Gene erfolgte dabei durch den bereits erwähnten RNAi-Mechanismus (s. vorne).

Zu einer ersten großen Serie von Experimenten mit durch RNAi stillgelegten Verteidigungsgenen des Wilden Tabaks gehörten transgene Pflanzen, in denen ein Schlüsselenzym des so genannten Jasmonat-Signalweges abgeschaltet wurde: die Lipoxxygenase. Dieses mit LOX bezeichnete Enzym vermittelt den Einstiegsschritt in die Biosynthese der Jasmonate (Abb. 5, Jasmonat ist das Salz der Jasmonsäure, die aus Membranlipiden freigesetzter Linolensäure gebildet wird). Jasmonate zählen zur Gruppe der Pflanzenhormone, wozu auch unter anderem das Salicylat, die Auxine und das gasförmige Reifungshormon Ethylen gehören. Das Auftreten von Jasmonat im Pflanzengewebe signalisiert sehr spezifisch den Fraß einer pflanzenfressenden Raupe, was dazu führt, dass die Pflanze den Verteidigungsfall auslöst – und dies innerhalb weniger Stunden nicht nur im angefressenen Blatt, sondern auch systemisch, also über die gesamte Pflanze verteilt. Wilder Tabak beginnt mit der Produktion von Nikotin und Protease-Inhibitoren (das sind Abwehrproteine, die die eiweißverdauenden Enzyme in der Raupe hemmen und der Raupe so den Appetit verderben; siehe auch Fallbeispiel 2) und synthetisiert Duftstoffe.

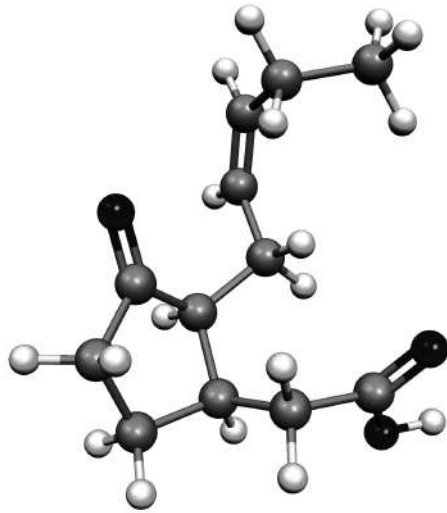


Abbildung 5: Molekül der Jasmonsäure, eines der wichtigsten Signalmoleküle in Pflanzen, das Abwehrreaktionen gegen Herbivoren aktiviert (Atome: weiß = Wasserstoff, grau = Kohlenstoff, schwarz = Sauerstoff). (Bildnachweis: MPI für chemische Ökologie).

Hintergrund der Versuche mit transgenen Pflanzen, denen es an Lipoxigenase-Aktivität mangelte, war zu prüfen, wie ein solcher an empfindlicher Stelle »entwaffneter« Wilder Tabak insbesondere im Freiland auf die Umwelt und auf Schädlingsbefall reagiert: Ist er nur noch bedingt abwehrbereit? Wie haushaltet er mit seinen Ressourcen? Wie ist es um seine *Darwinian Fitness* bestellt, das heißt: Wie viele fertile Samen produziert er im Vergleich mit anderen Versuchspflanzen, in denen die Lipoxigenase ganz normal arbeitet? Versuche mit im Gewächshaus kultivierten Pflanzen bestätigten, dass der transgene Tabak tatsächlich nicht mehr in der Lage war, seine Abwehr gegen einen seiner Hauptfeinde, den Pflanzenfresser *Manduca sexta*, in Gang zu setzen: Im Vergleich zu nicht transformierten Pflanzen (Wildtyp) gediehen die Raupen deutlich besser auf den transgenen asLOX (antisense-Lipoxigenase)-Pflanzen, messbar an ihrer jeweiligen Gewichtszunahme im gleichen Zeitraum. Chemische Analysen der Pflanzen zeigten, dass Nikotin und Protease-Inhibitoren (zur direkten Verteidigung) und der Duftstoff Bergamoten (zur indirekten Verteidigung)



Abbildung 6: Versuchsanordnung mit Pflanzen des Wilden Tabaks auf dem Versuchsfeld der *Lytle Ranch* in Utah, USA (Foto: André Kessler).

nur noch sehr vermindert gebildet wurden, und Jasmonat war in den Pflanzen, was zu erwarten war, nur noch in Spuren vorhanden [5]. Diese Ergebnisse bestätigten, dass durch das Abschalten von LOX die Pflanzen in ihrer Abwehrbereitschaft deutlich geschwächt waren. Somit war ein evidenter Beweis erbracht, dass Jasmonat ein entscheidender Signalgeber für das Vorliegen von Raupenfraß ist.

Überraschende Ergebnisse im Freilandversuch

Wie nun verhalten sich solche abwehrgeschwächten Tabakpflanzen im Freiland, und zwar an ihrem Ursprungsort im *Great Basin Desert* in Utah? Dort, wo sie nicht den wohligen Bedingungen eines Forschungsgewächshauses, sondern der rauen Natur mit ihren zahllosen abiotischen und biotischen Einflüssen ausgesetzt sind (Abb. 1)? Auf dem großen Versuchsfeld des Naturschutzgebietes *Lytle Ranch* wurden transgene Pflanzen mit nicht transgenen Wildtypfpflanzen gruppiert und in insgesamt 20 Gruppen randomisiert, also zufällig auf dem Feld verteilt freigesetzt (Abb. 6) [7]. Nach einer Woche der Anpassungsphase an ihre neue Umgebung wurden erste Untersuchungen an

den Pflanzen vorgenommen. Zunächst wurde festgestellt, dass sich alle Pflanzen im Hinblick auf ihre Wachstumsraten gleich verhielten. Auch in ihrer Erscheinung war kein Unterschied erkennbar: Alle Pflanzen waren im Rosettenblattstadium und verblieben gut entwickelt in der vegetativen Phase, das heißt es erfolgte während des gesamten Versuchszeitraums keine Blütenbildung. Bei Schädlingsbefall stellten sich Ergebnisse ein, die den im Gewächshaus in Jena erzielten vergleichbar waren: Der Lockstoff Bergamoten war in der Umgebung trotz Fraßinduktion bei den asLOX-Pflanzen nicht feststellbar, und *Manduca*-Raupe, die die Wissenschaftler auf den transgenen Pflanzen hatten fressen lassen, wiesen nach neun Tagen ein rund vierfach höheres Gewicht auf im Vergleich mit Raupen, die auf Wildtyp-Pflanzen gefressen hatten.

Weil sich die Tabakpflanzen auf dem freien Feld in Utah befanden und Tag und Nacht sämtlichen dort vorhandenen Feindseligkeiten ausgesetzt waren, konnten viele weitere, besonders ökologische Untersuchungen an den asLOX-Pflanzen erfolgen, die in einem *Containment* wie einem Gewächshaus eben nicht möglich sind. Das komplette Monitoring der Blattschädigungen, gemessen in Prozent der auf- oder angefressenen Blattflächen, ergab eine rund 7fach höhere Schädigung der asLOX-Pflanzen verglichen mit dem Wildtyp. Insgesamt 92 Prozent der asLOX- und 29 Prozent der Wildtyp-Pflanzen waren von Insekten attackiert worden, und zwar nicht präferenziell von ihren Hauptschädlingen wie *Manduca sp.* oder der Weichwanze *Tupiocoris notatus*, sondern von verschiedenen kauenden, stechenden, saugenden und minierenden Insektenarten. Die Analyse und Bestimmung dieser verschiedenen Herbivoren zeigte dann ein überraschendes Ergebnis: Zwei neue Arten traten auf dem Versuchsfeld auf, die in den vier vergangenen Jahren nicht oder nur in geringster Anzahl beobachtet worden waren und zudem nie einen nennenswerten Blattschaden hinterließen: eine Zwergzikaden-Art der Gattung *Empoasca* (Abb. 7) und der Blattkäfer *Diabrotica undecimpunctata tenella* Le Conte, ein Verwandter des gefährlichen Maiswurzelbohrers *Diabrotica virgifera virgifera*. Die *Empoasca*-Zikaden interessierten sich vor allem für die asLOX-Pflanzen: Diese wurden fast 15fach stärker geschädigt als der Wildtyp, 68 Prozent waren befallen (Wildtyp: 15 Prozent) und vor allem: Die Weibchen legten ihre Eier ausschließlich auf den asLOX-Pflanzen ab, aus denen junge Larven schlüpften, die den transgenen Tabak sofort als Babynahrung annahmen. Diese Beobachtung zeigt,



Abbildung 7: Eine Zwergzikaden-Art der Gattung *Empoasca*, die im Feld bevorzugt auf transgenem Wildem Tabak auftrat, bei dem ein Lipoxxygenase-Gen mittels RNAi stillgelegt worden war (Foto: André Kessler). (Farbabbildung 5, S. 146)

dass bereits die Induzierbarkeit pflanzlicher Abwehrmechanismen in einer gegebenen Pflanzenpopulation an einem gegebenen Standort das Spektrum der herbivoren, sie befallenden Insektenarten festlegt. Fällt ein Signalweg aus – wie in diesen Versuchen die durch Jasmonat vermittelte Induktion direkter und indirekter Verteidigung –, ändert sich deutlich sowohl die Zusammensetzung der Schädlingsarten als auch deren Populationsdichte. Welche Schädlinge eine Pflanze befallen, hängt also nicht nur davon ab, wie diese schmeckt (chemischer Phänotyp), sondern auch, auf welche Weise sich die Pflanze gegen ihre Feinde wehren kann [7].

Damit kommt ein neuer Aspekt chemischer Abwehrmechanismen von Pflanzen zum Vorschein: Pflanzen, die in ihrem natürlichen Lebensraum wachsen, sind unerbittlich nicht nur von ganz bestimmten, sondern von allen pflanzenfressenden Insekten bedroht. Doch wir nehmen viele dieser potenziellen Angreifer einfach deswegen nicht wahr, weil wir sie normalerweise nicht an der Pflanze fressend vor-

finden. Die mit Hilfe der transgenen Pflanzen gewonnen Ergebnisse zeigen nun, dass einige Herbivorenarten die Pflanzen anscheinend zuerst einmal darauf testen, ob sie als Nahrung geeignet sind. Reagiert die Pflanze auf diesen Test nicht mit einer chemischen Abwehr, wird sie sofort in die Speisekarte des Pflanzenfressers aufgenommen.

Fallbeispiel 2: Schwarzer Nachtschatten in Deutschland und den USA

Heimat und Lebensweise eines »Unkrauts«



Abbildung 8: Wildtyp (nicht gentechnisch veränderte Pflanze) des Schwarzen Nachtschattens (*Solanum nigrum*) auf einem Versuchsfeld nahe Dornburg in Thüringen, Deutschland (Foto: Markus Hartl). (Farbabbildung 3, S. 145)

Der Schwarze Nachtschatten (*Solanum nigrum*) ist eine Pflanze, die wohl viele als Unkraut bezeichnen würden (Abb. 8). Sein Ursprung liegt wahrscheinlich in Eurasien oder Afrika, mittlerweile hat er sich aber auf der ganzen Welt ausgebreitet bzw. wurde vom Menschen verschleppt. Er ist wie der Tabak, die Tomate oder die Kartoffel ein Vertreter der Nachtschattengewächse (Solanaceae) und wächst vorzugsweise in nährstoffreichen, gestörten Habitaten (z.B. Flußauen) und auf Ruderalflächen, also vom Menschen geprägten Standorten wie Brachflächen, Deponien oder Bahnhöfen, und kommt auch als Ackerbeikraut vor. Der direkte Nutzen für den Menschen ist begrenzt, da die Pflanze giftige Steroidalkaloide produziert, die sich in Blättern und Früchten ansammeln.⁵

- ⁵ Das nach dem Nachtschatten benannte Solanin kommt übrigens nebst anderen Alkaloiden auch in grünen Tomaten oder in ergrünenden Kartoffelknollen vor, weshalb von deren Verzehr abgeraten wird.

Dennoch wird die Pflanze in einigen Ländern als Spinat zubereitet oder die reifen Früchte werden als Obst verzehrt, wobei es sich dabei um Varietäten mit niedrigerem Alkaloidgehalt handeln könnte. Es liegt auf der Hand, dass die giftigen Alkaloide, ähnlich dem Nikotin im Tabak, die Pflanzen wahrscheinlich vor Fraßfeinden schützen. Wie sich der Wilde Tabak nicht nur auf das Nikotin zur Abwehr verlässt, so hat auch der Schwarze Nachtschatten noch andere Möglichkeiten, sich zur Wehr zu setzen. Ähnlich dem Tabak ist er in der Lage, eine Verletzung durch einen Herbivoren zu erkennen und daraufhin selektiv über das Hormon Jasmonsäure bestimmte Gene zur Verteidigung zu aktivieren; man spricht in diesem Fall von induzierter Abwehr. Beim Schwarzen Nachtschatten haben uns allerdings nicht wie im vorigen Beispiel die indirekte Abwehr durch Duftstoffe interessiert, sondern die direkten Abwehrsysteme, also Substanzen, die eine direkte negative oder abschreckende Wirkung auf die Herbivoren haben.

Das Verursachen von Bauchschmerzen als Abwehrstrategie

Eines der bekanntesten Beispiele einer derartigen induzierten, direkten Abwehr wurde zum ersten Mal in Tomaten beschrieben. Dabei wurde beobachtet, dass Tomatenblätter nach Verletzung eine bestimmte Art von Proteinen produzieren, die die Eigenschaft besitzen, die Verdauungsenzyme (Proteasen) ihrer Fraßfeinde zu hemmen, weswegen sie den Namen Protease-Inhibitoren erhalten haben. Mittlerweile wurden zahlreiche verschiedene Protease-Inhibitoren aus einer Vielzahl von Pflanzenarten isoliert und beschrieben. In einigen Pflanzen werden sie nur nach Verwundung, in manchen besonders wichtigen Geweben wie z.B. den Samen dauerhaft produziert, wahrscheinlich um diese noch besser zu schützen.⁶ Viele der Protease-Inhibitoren haben tatsächlich eine starke Hemmwirkung auf Verdauungsenzyme von Insekten, aber auch von Säugetieren und bereiten ihnen beträchtliche Bauchschmerzen. So sind rohe Kartoffeln unter anderem deswegen so schwer ver-

- 6 Der Grund, warum die Pflanze die Protease-Inhibitoren nicht in allen Geweben ständig produziert, ist wahrscheinlich der, dass es ziemlich kostspielig wäre. Ganz allgemein wird angenommen, dass der Vorteil induzierter Abwehrsysteme darin liegt, dass sie tatsächlich nur dann aktiviert werden, wenn sie gebraucht werden, also wenn die Pflanze tatsächlich attackiert wird. In »Friedenszeiten« kann die dadurch gesparte Energie in Wachstum, Entwicklung und Vermehrung investiert werden.

daulich, weil bis zu 50 % ihres löslichen Proteins aus derartigen Inhibitoren besteht, die erst beim Kochen inaktiviert werden. Mittlerweile hat man sehr viel über Protease-Inhibitoren gelernt: Man hat sie isoliert und ihre Hemmwirkung auf eine Vielzahl von Proteasen getestet, man hat manche der Inhibitoren mit gentechnischen Methoden in andere Pflanzen übertragen, um in Laborversuchen zu erforschen, ob diese besser gegen Insekten geschützt sind, man hat Insekten mit Inhibitoren getränktes Futter gegeben, um zu sehen, wie sie darauf reagieren, und vieles mehr. Zunächst war man darauf konzentriert, die allgemeinen Eigenschaften und Wirkmechanismen zu erforschen und verwendete dafür Modellsysteme, die entweder gut für Laborversuche geeignet oder die aus pflanzenzüchterischer Sicht von Interesse waren (typische Nutzpflanzen und deren Schädlinge). Dabei kamen allerdings grundlegende ökologischen Fragestellungen zu kurz: 1. Funktionalisieren Protease-Inhibitoren tatsächlich als Abwehr gegen die natürlichen Herbivoren einer Pflanze? 2. Müssten sich die typischen Fraßfeinde einer Pflanzenart nicht schon längst an eine derartige Abwehr angepasst haben, gibt es also Unterschiede in der Wirkung auf verschieden angepasste Insekten? 3. Haben die Protease-Inhibitoren noch andere Funktionen, die nichts mit der Abwehr von Herbivoren zu tun haben?

Die Protease-Inhibitoren des Schwarzen Nachtschattens

Es zeigte sich, dass der Schwarze Nachtschatten zur Untersuchung dieser Frage wie geschaffen scheint. In unseren Experimenten konnten wir herausfinden, dass er nicht nur einen, sondern vier verschiedene Protease-Inhibitoren produziert [8]. Allen vier ist gemeinsam, dass sie nach Verwundung und Insektenfraß verstärkt gebildet werden, also induzierbar sind – allerdings in unterschiedlichem Ausmaß. Andererseits gibt es wichtige Unterschiede zwischen den vier Protease-Inhibitoren: Sie werden in verschiedenen Geweben der Pflanze unterschiedlich stark produziert – zum Teil auch ohne Insektenbefall – und sie zeigen beachtliche Unterschiede in ihrem Hemmpotenzial gegenüber verschiedenen Proteasen. Das ließ vermuten, dass die vier zugrunde liegenden Gene unter Umständen auch verschiedene Funktionen haben könnten. Um dieser Vermutung näher auf den Grund zu gehen, haben wir mit der oben beschriebenen RNAi-Methode transgene Pflanzen hergestellt, in denen die vier Gene einzeln oder in

Kombination ausgeschaltet waren. In Experimenten im Gewächshaus zeigte sich dann allerdings ein überraschendes Ergebnis. Lässt man die bereits erwähnten Raupen von *Manduca sexta* an diesen transgenen Pflanzen fressen, so unterscheiden sie sich in ihrem Wachstum nicht von Raupen, die ausschließlich an Wildtyp-Pflanzen fressen. Die vier Inhibitoren sind als Verteidigung gegen dieses Insekt also völlig wirkungslos. Dabei hätten wir erwartet, dass das Fehlen der Verdauungshemmer in den transgenen Pflanzen einen Vorteil für die Insekten bringt, weil sie die Nahrung besser verdauen und umsetzen können sollten. Um herauszufinden, ob es sich dabei um ein allgemeines Phänomen handelt, haben wir ähnliche Fraßversuche auch mit anderen Herbivoren durchgeführt: Larven der Baumwoll-Eule (*Spodoptera littoralis*) und der Zuckerrüben-Eule (*Spodoptera exigua*). Beide sind Vertreter der Eulenfalter (Noctuidae) und bekannt dafür, an einer großen Zahl von Pflanzen inklusive *Solanum nigrum* zu fressen und beträchtlichen (wirtschaftlichen) Schaden anzurichten. Interessanterweise haben sich die beiden Arten in den Fraßversuchen eindeutig unterschieden. Während die Zuckerrüben-Eule auf Pflanzen, bei denen zwei oder mehr Protease-Inhibitoren ausgeschaltet waren, an Gewicht zunahm, schien die Baumwoll-Eule von der An- oder Abwesenheit der Inhibitoren ähnlich wie *Manduca sexta* nicht berührt zu werden. In den Gewächshausversuchen erwiesen sich die Protease-Inhibitoren also als nur mäßig erfolgreiche Abwehr gegen Herbivoren, stark abhängig von der Schädlingsart [8].

Nun waren wir natürlich daran interessiert herauszufinden, wie sich die Pflanzen ohne Protease-Inhibitoren in ihrem natürlichen Lebensraum behaupten würden. Auf einer Versuchsfläche nahe Dornburg bei Jena haben wir dazu die transgenen Pflanzen und Wildtyppflanzen ausgebracht und die auftretenden Herbivoren und den von ihnen verursachten Schaden protokolliert. Dabei zeigte sich, dass die Pflanzen vor allem von einer Flohkäferart (*Epitrix pubescens*, Abb. 9) befallen wurden, einem absoluten Spezialisten für Nachtschattengewächse. Der Flohkäfer ist, wie der Name vermuten lässt, zwar sehr klein, durch sein massenartiges Auftreten innerhalb kurzer Zeit kann er aber beträchtlichen Schaden an den Blättern des Schwarzen Nachtschattens anrichten. Bisher war es uns nicht gelungen, den Flohkäfer in stabilen Kulturen im Labor zu züchten, weswegen das natürliche Auftreten der Käfer auf dem Versuchsfeld eine einmalige Gelegenheit darstellte. Das Ausschalten der Protease-Inhibitoren hatte allerdings, ähnlich wie bei



Abbildung 9: Flohkäfer (*Epitrix pubescens*) auf einem Blatt des Schwarzen Nachtschattens mit den für den Käfer charakteristischen kreisrunden Fraßstellen (Foto: Danny Kessler). (Farbabbildung 6, S. 147)

den Raupenversuchen im Gewächshaus, keinen nennenswerten Einfluss auf das Verhalten der Flohkäfer. Transgene und Wildtyp-Pflanzen wurden etwa gleich stark befallen. Sollten die Protease-Inhibitoren am Ende zur Verteidigung gar nicht geeignet sein?

Zur Klärung dieser Frage beschlossen wir, einen weiteren Versuch zu unternehmen und beantragten eine Genehmigung, die entsprechenden transgenen Pflanzen auch auf der *Lytle Ranch* in den USA aussetzen zu dürfen – eben dort, wo wir auch unsere Versuche mit dem Wilden Tabak durchführen. Die Idee dahinter war, dass der Schwarze Nachtschatten erst im Laufe der letzten Jahrhunderte in die USA eingeschleppt wurde und daher eine evolutionär eher kurze Zeitspanne für eine wechselseitige Anpassung von Pflanze und herbivoren Insekten zur Verfügung stand. Wir erwarteten also, dass die Insekten in Utah vielleicht anders auf die Protease-Inhibitoren des Schwarzen

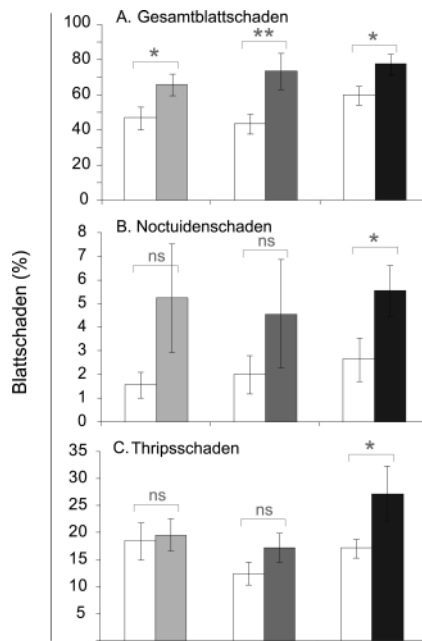


Abbildung 10: Durch Insektenbefall verursachter Schaden an Blättern des Schwarzen Nachtschattens in Feldversuchen in Utah. A: durchschnittlicher Gesamt-Blattschaden pro Pflanze. B: durch Eulenfalter-Larven verursachter durchschnittlicher Blattschaden. C: durch Thripse verursachter durchschnittlicher Blattschaden (alle Angaben in Prozent). Transgene Pflanzen, bei denen eine unterschiedliche Anzahl von Protease-Inhibitoren ausgeschaltet worden war, wurden in Pärchen mit unveränderten Wildtyp-Pflanzen ausgepflanzt und der auftretende Schaden protokolliert. Balkenfarben: weiß = Wildtyp, grau = Pflanzen, bei denen zwei (hellgrau), drei (mittelgrau) oder alle vier (dunkelgrau) Protease-Inhibitoren gentechnisch ausgeschaltet wurden. * bezeichnet statistisch signifikante Unterschiede (* < 0,5 %, ** < 0,1 % Fehlerwahrscheinlichkeit), ns = nicht signifikant (aus [8]).

Nachtschattens reagieren, weil sie noch nicht so gut an sie angepasst sind. Tatsächlich zeigte sich in diesem Feldversuch ein anderes Bild als in Deutschland. Alle transgenen Linien, egal ob zwei, drei oder vier Protease-Inhibitoren ausgeschaltet waren, zeigten einen etwa um 20 bis 30 % höheren Gesamt-Blattschaden als der Wildtyp (Abb. 10). Das bedeutet, dass am Standort Utah die Protease-Inhibitoren für den

Schwarzen Nachtschatten durchaus eine wirksame Verteidigung gegen Herbivoren darstellen [8]. Für eine genauere Analyse haben wir den beobachteten Schaden nach den verschiedenen Insektengruppen unterteilt und daraus weitere interessante Schlüsse ziehen können, von denen wir zwei Beispiele herausgreifen wollen. Die Raupen der in Utah vorkommenden Vertreter der Eulenfalter scheinen erst durch die Kombination aller vier Inhibitoren signifikant in ihrem Fressverhalten beeinflusst zu werden. In den transgenen Linien mit nur zwei bzw. drei ausgeschalteten Inhibitoren zeigte sich zwar auch ein durchschnittlich höherer Schaden als im Wildtyp, der aber statistisch gesehen nicht signifikant war. Das bedeutet, der beobachtete Unterschied bei diesen Linien könnte entweder rein zufällig gewesen sein oder aber, wenn es ein tatsächlicher Unterschied gewesen ist, so war die Schwankung der Werte innerhalb der Beobachtungen so groß, dass es weiterer und unter Umständen größer angelegter Versuche bedurfte, um dies zweifelsfrei festzustellen. Anders sah das Bild bei einer weiteren Herbivorenklasse aus, den Thripsen (Fransenflügler, Thysanoptera). Diese sehr kleinen Insekten, die sich vor allem durch Anstechen und Aussaugen einzelner Blattzellen ernähren, reagierten offenbar auf das Ausschalten von drei der Protease-Inhibitoren überhaupt nicht. Erst nach Ausschalten des vierten zeigte sich ein 10 % höherer Schaden als am Wildtyp (Abb. 10).

Wie sind all diese Ergebnisse nun zu bewerten? Die Feldversuche in Utah beweisen, dass Protease-Inhibitoren tatsächlich in der Lage sind, den Schwarzen Nachtschatten gegen bestimmte Insektenarten zu verteidigen – allerdings haben die vier verschiedenen Inhibitoren unterschiedlich starke Wirkung auf verschiedene Herbivorengruppen. Die Feld- und Laborversuche in Deutschland zeigten aber auch, dass manche Insektenarten gegen diese Protease-Inhibitoren resistent zu sein scheinen. Was ist nun der Unterschied zwischen diesen resistenten Insekten und den nicht-resistenten? Sowohl Flohkäfer als auch *Manduca sexta* sind ausgeprägte Nachtschattenspezialisten. In der co-evolutionären Anpassung an ihre Futterpflanzen haben sie offenbar Strategien entwickelt, mit deren Abwehrmechanismen umzugehen. Von *Manduca sexta*-Raupen ist bekannt, dass sie in der Lage sind, mit den Protease-Inhibitoren des Wilden Tabaks zurechtzukommen, die in ihrer Konzentration weit über denen des Schwarzen Nachtschattens liegen. Es überrascht daher nicht, dass das Ausschalten der Inhibitoren des Schwarzen Nachtschattens keine weitere Wirkung nach sich zieht.

Der Wilde Tabak kann sich gegen *Manduca sexta* nur durch eine Kombination aus Protease-Inhibitoren, Nikotin und anderen Abwehrmaßnahmen erfolgreich zur Wehr setzen und es ist zu erwarten, dass auch beim Schwarzen Nachtschatten andere Abwehrmechanismen gegen bestimmte Insekten von größerer Bedeutung sind als die Protease-Inhibitoren (wie z.B. die erwähnten Steroid-Alkaloide). Der Flohkäfer hingegen könnte eine andere Strategie verfolgen: Da er sich immer nur für kurze Zeit an einer Pflanze aufhält, könnte es sein, dass er bereits zur nächsten Pflanze weitergezogen ist, bevor die Protease-Inhibitoren überhaupt in nennenswerten Mengen im Blatt gebildet wurden, was in der Regel zwei bis drei Tage dauert. Der Flohkäfer reagiert damit schneller als die Abwehrsysteme und würde sie so ganz einfach vermeiden (was aber experimentell erst noch zu beweisen wäre). Im Gegensatz zu den beiden Spezialisten reagieren eine Reihe von Herbivoren durchaus sensibel auf die Protease-Inhibitoren. Bei diesen Herbivoren handelt es sich in der Regel um weniger spezialisierte Insekten, die auf verschiedenen Pflanzenarten fressen. Da der Schwarze Nachtschatten meist in geringer Zahl an gestörten Standorten auftritt, also keine über mehrere Jahre stabilen, großen Populationen bildet, ist die Wahrscheinlichkeit, auf eher unspezialisierte Herbivoren zu treffen, relativ groß. Gegen diese so genannten Generalisten sind die Protease-Inhibitoren des Schwarzen Nachtschattens eine wirksame Verteidigung. Allerdings gilt in der Biologie oft »keine Regel ohne Ausnahme«: Die Baumwoll-Eule ist keineswegs auf Nachtschattengewächse spezialisiert, ließ sich aber von den Protease-Inhibitoren des Nachtschattens trotzdem nicht beeindrucken. Bei dieser Art handelt es sich aber tatsächlich um einen besonderen Fall. Die Baumwoll-Eule gilt als extrem polyphag, das bedeutet, sie frisst an einer sehr großen Anzahl von verschiedenen Pflanzenarten, darunter vielen Nutzpflanzen, und zählt deshalb zu den weltweit zerstörerischsten Schadinsekten im Pflanzenbau. Wie die Baumwoll-Eule es schafft, die Abwehrsysteme so vieler verschiedener Pflanzenarten zu überwinden, ist ein Rätsel für sich und Gegenstand intensiver Forschung.

Die Labor- und vor allem die Feldversuche mit dem Schwarzen Nachtschatten haben demonstriert, dass Protease-Inhibitoren durchaus eine effiziente Abwehr gegen Insekten darstellen können. Allerdings wirken bestimmte Inhibitoren nur gegen bestimmte Insektenarten und zum Teil auch nur in Kombination miteinander. Stark spezialisierte Herbivoren, von denen der Flohkäfer ausschließlich im Feld

untersucht werden konnte, haben sich aber offenbar so gut an diese Abwehr angepasst, dass die Protease-Inhibitoren alleine in diesen Fällen als Verteidigung wirkungslos sind. Das könnte auch erklären, warum der Schwarze Nachtschatten über mehrere Inhibitoren mit verschiedenen Wirkungen verfügt. Die Vielfalt könnte ein Ausdruck der vielfältigen Anforderungen sein, die durch zahlreiche Insekten an die Pflanze gestellt werden. Abgesehen davon war allerdings noch eine Beobachtung bemerkenswert: Zwei der vier Protease-Inhibitoren hemmen bevorzugt Verdauungsenzyme, die in Insekten eher selten sind. Weitaus häufiger kommt diese Klasse in Mikroorganismen oder in den Pflanzen selbst vor. Das könnte bedeuten, dass die Inhibitoren neben der Verteidigung gegen Insekten auch Funktionen in der Abwehr pathogener Mikroorganismen haben könnten oder den Abbau von pflanzeigenen Proteinen durch pflanzliche Proteasen kontrollieren. Das würde eine weitere Erklärung liefern, warum nicht alle Protease-Inhibitoren gleich stark gegen Insekten wirksam sind und würde außerdem bedeuten, dass diese Gene vielleicht mehrere Funktionen zu erfüllen haben. Noch ist das Spekulation, aber wir erwarten mit Spannung die Durchführung weiterer Versuche, die diesen Fragen auf den Grund gehen werden.

Zusammenfassung

Wir hoffen, dass wir einen kleinen Einblick gewähren konnten, welche Herausforderungen die Ergebnisse der Genomforschung für die gesamten Biowissenschaften bereithalten. Auf der Ebene des Genoms zu verstehen, wie ein Organismus wächst, überlebt und sich fortpflanzt, ist eine gleichermaßen spannende wie schwierige Aufgabe. Vom erfolgreichen Bestehen dieser Herausforderungen werden zu einem großen Teil unsere zukünftigen Fortschritte nicht nur in der Medizin, sondern in allen Bereichen der Lebenswissenschaften abhängen. Dazu zählt insbesondere auch die ökologische Forschung, die uns die Grundlagen dafür bieten kann, zu verstehen, wie Ökosysteme funktionieren, was zu ihrem Erhalt notwendig ist und wie sie durch das menschliche Handeln beeinflusst werden. Die Verbindung von molekularen Methoden mit den Fragestellungen und Methoden der Ökologie ist daher eine logische und notwendige Konsequenz, um zukünftigen Aufgaben und neuen Fragestellungen angemessen zu begegnen.

Grüne Gentechnik ist weltweit seit den 1980er Jahren in den Pflanzenwissenschaften eine alltägliche Angelegenheit. Sie gehört zu den gängigen Methoden, die jeder Wissenschaftler anwenden kann, um beispielsweise die für die Pflanze lebensnotwendigen Stoffwechsel- und Regulationsvorgänge zu entschlüsseln und erklären zu können. Mutanten oder transgene Linien mit veränderter Genaktivität in einem gegebenen Versuchsansatz sind der *gold standard of proof*, also ein wissenschaftlich akzeptierter Beweis. Der inzwischen in Lehrbüchern beschriebene Zuckerstoffwechsel von Pflanzen mit allen seinen Komponenten – von der Photosynthese über Transportvorgänge in den Leitgeweben bis zur Stärkeproduktion – wurde nicht unwesentlich mithilfe transgener Pflanzen aufgeklärt. Um die Biochemie von Stoffwechselwegen zu studieren, reichen Experimente im Gewächshaus, in Klimakammern oder im Labor aus, und für bestimmte Analysen sind kontrollierte Umweltbedingungen, wie man sie in einer Klimakammer einstellen kann, sogar dringend erforderlich. Um aber das Stoffwechselgeschehen einer Pflanze im Tages- und Jahreszeitenverlauf an einem gegebenen Standort im Zusammenhang mit den dort vorkommenden weiteren Pflanzen- und Tierarten zu verstehen, um also ökologische Forschung zu betreiben, sind Freilandversuche unabdingbar. Freisetzen genetisch veränderter Pflanzen – ob natürliche Mutanten oder gentechnisch veränderte Pflanzen – sind erforderlich, wenn mit ihrer Hilfe bestimmte Merkmale wie beispielsweise die der natürlichen, genetisch verankerten Schädlingsresistenz genau studiert werden können. Die beschriebenen Fallbeispiele Wilder Tabak und Schwarzer Nachtschatten bestätigen und erhärten diese Auffassung.

Unsere Experimente zeigen zudem den wissenschaftlichen Wert natürlicher Lebensräume und ihrer Vielfalt. Die Natur war und ist *per se* das Labor für Naturforscher und Ökologen, in dem die Experimente der Evolution ihren Lauf nehmen und deren Ergebnisse in den Genomen der Lebewesen festgehalten werden. Kein noch so modern ausgerüstetes Labor kann diese natürlichen Gegebenheiten in allen Teilen nachbilden, kein Computermodell kann sie simulieren – dazu wissen wir, zumindest im Moment, einfach noch zu wenig. Solange nicht jeder Bestandteil und jede Beziehung innerhalb eines Ökosystems bekannt sind, kann ökologische Grundlagenforschung im Labor alleine nicht funktionieren. Die Vielfalt der Natur ist daher unser Lehrmeister, der uns die ökologischen Zusammenhänge und die einzelnen Bestandteile von Ökosystemen näherbringt – bis wir letzt-

lich vielleicht auch begreifen, wie Ressourcen schonend und nachhaltig zu nutzen sind. Bis es soweit ist, gibt es noch unglaublich viel zu lernen und zu entdecken. Ein verantwortungsvoller Einsatz von gentechnisch veränderten Wildpflanzen kann auf diesem Weg eine wichtige Hilfe sein.

Ein Kommentar: Genehmigung, Durchführung, Sicherheit und Akzeptanz gentechnisch-ökologischer Freilandversuche in den USA und Deutschland

Der große Diskussionsbedarf, der beim *XLAB Science Festival* nach dem Vortrag entstanden ist, hat gezeigt, wie viel Interesse an dem Thema Grüne Gentechnik vor allem im Zusammenhang mit möglichen Risiken besteht. Aus diesem Grund wollen wir an dieser Stelle unseren persönlichen Standpunkt zu diesem Thema verdeutlichen. In der Grundlagenforschung werden zu untersuchende Gene einer Pflanze meistens ausgeschaltet oder verstärkt, um nachfolgend im Vergleich mit nicht veränderten Pflanzen das durch das jeweilige Gen ausgeprägte Merkmal zu studieren. Bei dem dafür erforderlichen gentechnischen Eingriff ist die Menge artfremden Genmaterials auf ein Minimum beschränkt, z.B. auf Selektionsmarker zur Auswahl der erfolgreich veränderten Pflanzen. Das Risiko einer derartigen genetischen Veränderung strebt für Mensch und Umwelt gegen null, wie seit vielen Jahren zahlreiche unabhängige Behörden und Wissenschaftler weltweit bestätigen konnten, denn es handelt sich größtenteils um genetische Veränderungen, die in der Natur jederzeit und vor allem unkontrolliert vorkommen. Durch spontane, natürliche Mutationen werden in der Natur in sämtlichen Lebensformen, auch beim Menschen, kontinuierlich Genaktivitäten inaktiviert oder verstärkt und mehr oder weniger große DNA-Abschnitte im Genom hin und her bewegt. Falls dies für das Überleben oder die Fortpflanzung des jeweiligen »mutierenden« Organismus von Nachteil ist, wird eine solche genetische Veränderung durch natürliche Selektion schnell aus der Population verschwinden. Genauso verhält es sich mit einer durch den Menschen erzeugten transgenen Pflanze, bei der ein Gen zu Versuchszwecken ausgeschaltet wurde: Der Verlust stellt in der Regel einen Nachteil dar. Das bedeutet, dass bei einer unkontrollierten, über die Versuchsfelder

hinausgehenden Verbreitung einer solchen Pflanze das transgene Merkmal schon nach wenigen Jahren vermutlich nicht mehr aufzufinden ist. Das ist umso wahrscheinlicher, wenn es sich bei den untersuchten Genen wie in unserem Beispiel um Vertreter handelt, die für die Abwehr der Pflanze gegen Schädlinge essenziell sind.

Ist ein solches Gen und seine Funktion für die Abwehr bestimmter Schädlinge erkannt und seine Rolle und Wirkung durch gentechnisch-ökologische Freilandexperimente belegt, so kann dies eine wertvolle Information für Pflanzenzüchter darstellen. Im Zuge moderner Züchtungsverfahren wie der Präzisionszucht, dem so genannten *SMART Breeding* (*Selection with Markers and Advanced Reproductive Technologies*), werden Sequenzen oder chromosomale Positionen eines Resistenzgens als molekulare Marker verwendet, um die aus konventioneller Züchtung hervorgegangenen Nachkommen und vorab deren Eltern auf dieses Merkmal hin zu testen. Weitergezüchtet werden dann nur noch diejenigen Nachkommen, bei denen sich das Resistenzgen nicht »herausgemendelt« hat. Wohlgemerkt: Es handelt sich hier nicht um die Erzeugung neuartiger transgener Nutzpflanzen, sondern um die Optimierung und Beschleunigung der konventionellen Pflanzenzüchtung mittels molekularer Marker. Die Grüne Gentechnik kann somit sowohl durch die Erzeugung transgener Nutzpflanzen, die beispielsweise insekten- oder pilzresistent sind, als auch durch ihren mittelbaren Einsatz in den Pflanzen- und Agrarwissenschaften einen Beitrag zu einer fortschrittlichen Agrikultur leisten, die angesichts abnehmender Anbauflächen bei gleichzeitigem Anstieg der Weltbevölkerung dringend erforderlich scheint.

Wir möchten an dieser Stelle auch darauf hinweisen, dass die Studien am MPI für chemische Ökologie weder von der Industrie gefördert sind noch Kooperationen bestehen, die Daten für Agrar- oder andere Firmen direkt nutzbar machen. Natürlich sind manche unserer Ergebnisse, die öffentlich in Fachzeitschriften publiziert werden, für Pflanzenzüchter von Interesse, aber unsere Forschung verläuft nicht zielorientiert in Richtung einer bestimmten Anwendung – dies würde sich auch nicht mit der primären Aufgabe der Max-Planck-Gesellschaft, nämlich Grundlagenforschung zu betreiben, vertragen.

Die Behörden, mit denen wir im Laufe unserer Freisetzungversuche zusammengearbeitet haben, haben wir als sehr gewissenhafte und auf vollständige Sicherheit bedachte Kontrollorgane erlebt. Um Art und Umfang der Sicherheitskontrollen etwas näher zu bringen,

wollen wir einen kleinen Einblick in den Genehmigungsprozess für Freisetzungsversuche diesseits und jenseits des Atlantiks gewähren. Die durch das Max-Planck-Institut für chemische Ökologie erfolgten Freisetzungen gentechnisch veränderter Wildpflanzen in den USA (seit 2001) und in Deutschland (2004 bis 2007) haben uns viele Erfahrungen im Umgang mit dem Thema Grüne Gentechnik sammeln lassen. In Europa und in den USA bedarf es bürokratischen als auch experimentellen Aufwands, um eine Genehmigung für Freisetzungen zu bekommen. Der bürokratische Aufwand ist in den USA geringer, und ein Antrag wird in der Regel innerhalb von vier Wochen bearbeitet. In Deutschland liegt die durchschnittliche Bearbeitungszeit bei rund sechs Monaten. Ein Grund dafür ist, dass der Antrag neben der Genehmigungsbehörde drei weiteren Behörden, unter anderem dem Bundesamt für Naturschutz, vorgelegt werden muss. Experimenteller Aufwand entsteht, weil freizusetzende Pflanzen beispielsweise auf Anzahl und Vollständigkeit von T-DNA-Insertionen oder auf die Anwesenheit von Vektorsequenzen überprüft werden müssen. Hinzu können zahlreiche weitere Daten zur Beschreibung des jeweiligen Transgens, zum Grad des *Gene Silencings* eines stillzulegenden Gens sowie zur Beschreibung aller im Gewächshaus auftretenden Unterschiede in Wachstum und Metabolismus der transgenen Pflanze im Vergleich mit dem Wildtyp kommen. Die folgenden Angaben betreffen vornehmlich Freisetzungsanträge in Deutschland: Vorgenommen werden muss eine Abschätzung von potenziellen Wechselwirkungen zwischen den Versuchspflanzen und Ziel- bzw. Nichtzielorganismen und deren mögliche Folgen für Mensch und Umwelt. Weiterhin sind Angaben erforderlich über das Vorhandensein geschützter Arten um das Versuchsfeld herum (Pflanzen- und Tierarten), und auch der Abstand zu offiziell anerkannten Naturschutzgebieten muss angegeben werden. Alle diese Daten zusammen mit Angaben zur Naturgeschichte der gentechnisch veränderten Pflanzenart sowie ihren Verbreitungsgebieten besonders in der Umgebung des Versuchsfeldes, Aufzählungen verwandter Arten und ihrer Verbreitung und nicht zuletzt ausführliche Angaben zu Ablauf und Art der genetischen Veränderung bilden die Grundlage für eine Risikoabschätzung und Umweltverträglichkeitsprüfung der transgenen Pflanze. Im Antrag enthalten ist weiterhin ein Überwachungsplan, der festlegt, wann, wie oft und was an den transgenen Pflanzen, sobald sie im Feld sind, kontrolliert, vorgenommen und protokolliert werden muss – neben den eigentlichen experi-

mentell bedingten Messungen. Die Genehmigungsbehörde prüft den Antrag auf Vollständigkeit, widmet sich der Risikoabschätzung und der Umweltverträglichkeitsprüfung und ergänzt gegebenenfalls den Überwachungsplan durch die Aufstellung verpflichtender Nebenbestimmungen.

Die jeweils zuständigen Genehmigungsbehörden, in den USA der *Animal and Plant Health Inspection Service (APHIS)* des *US Departments of Agriculture* und in Deutschland das Bundesamt für Verbraucherschutz und Lebensmittelsicherheit (BVL), erscheinen uns in ihrer Genehmigungspraxis politisch, also bedingt durch die jeweils in beiden Staaten öffentlich vertretenen Meinungen zum Thema Grüne Gentechnik, die über die Parteien Einzug in die Legislative gefunden haben, unterschiedlich geprägt. Den Unterschied kann man wie folgt auf einen Punkt verdichten: »*Transgenic plants are safe – let's prove them unsafe*« – dies entspricht in etwa der Genehmigungspraxis in den USA. »*Transgenic plants are unsafe – let's prove them safe*« – dies stellt die Meinung in Deutschland und in weiten Teilen Europas dar, mit der sich das Bundesamt für Verbraucherschutz und Lebensmittelsicherheit in seiner Funktion als Bundesoberbehörde auf der Grundlage der nationalen Gesetzgebung und europäischer Richtlinien auseinandersetzen muss. Der deutsche, pessimistische Ansatz, dass transgene Pflanzen unabhängig von ihrer jeweiligen genetischen Veränderung als Risiko betrachtet werden, schlägt sich nicht selten schon in einer eher zurückhaltend formulierten Antragstellung seitens der Pflanzenforscher nieder. Um einer Genehmigung möglichst wenig in den Weg zu stellen, werden geplante Freisetzungsversuche auf ein Minimum reduziert (kleine Anzahl freizusetzender Pflanzen, sichere Standorte z.B. unmittelbar auf dem Gelände des Forschungsinstituts). Dennoch erfolgen im Zuge einer Genehmigung oft zusätzliche behördliche Auflagen. Während unserer Freisetzungen von transgenem Schwarzen Nachtschatten in Deutschland musste beispielsweise das Versuchsfeld im Umkreis von 35 Metern alle vier Wochen auf sieben verwandte Arten der Gattung *Solanum* durchsucht und diese gegebenenfalls entfernt werden, um einem potenziellen Auskreuzen des transgenen Merkmals vorzubeugen, obwohl beantragt worden war, in allen Experimenten ohnehin nur nichtblühenden Schwarzen Nachtschatten einzusetzen, der also keinen Pollen in die Umwelt abgeben kann. Im Überwachungsplan hatten wir zusätzlich manifestiert, dass alle zwei Tage sämtliche Pflanzen auf das Auftreten von Blütenknospen überprüft und diese voll-

ständig entfernt werden, sobald sie sichtbar sind (Abb. 11). Das Entfernen der Knospen wurde je Pflanze dokumentiert und alle Aufzeichnungen wurden jährlich von der thüringischen Überwachungsbehörde überprüft. Obwohl nach unserer Einschätzung dank aller unserer Vorkehrungen und der botanischen Erkenntnisse über *Solanum nigrum* ein Auskreuzen ausgeschlossen werden konnte, wurde das Absuchen und Entfernen von sieben weiteren *Solanum*-Arten angeordnet, wobei fraglich bleibt, ob ein zwischenartliches Kreuzen zweier *Solanum*-Arten überhaupt zu fertilen Nachkommen geführt hätte.

Insgesamt leistet sich Deutschland wegen seiner sehr kritischen Einstellung gegenüber der Grünen Gentechnik schon seit Beginn des neuen Jahrtausends eine millionenschwere Biosicherheitsforschung. Mit ihr wird versucht, empirisch alle möglichen Risiken für Umwelt und Mensch sogar von seit Jahren auf Millionen von Hektar angebauten transgenen Nutzpflanzen zu ermitteln, wie beispielsweise des Insektenfraß-resistenten Mais MON810 (»Bt-Mais«). Die Ergebnisse aus der Sicherheitsforschung haben zwar Bedenken gegen transgene Pflanzen aus dem Weg räumen können. Jedoch wird die Biosicherheitsforschung insbesondere von Seiten der radikalen Gentechnik-kritiker nicht akzeptiert und damit einhergehend werden alle, auch die nur unter besonderen Auflagen und nach gründlicher Prüfung durch die Zentrale Kommission für Biologische Sicherheit amtlich zugelassenen, Freisetzungsexperimente kategorisch abgelehnt und Versuchsfelder zum größten Teil attackiert und zerstört. »Es gibt nichts Gutes, außer man tut *nichts*«, so könnte man Erich Kästners Aussage aus der Sicht der Gentechnikgegner abwandeln.

Die praktische Durchführung gentechnisch-ökologischer Freilandversuche, so soll am Schluss kritisch zusammengefasst werden, ist in Deutschland aufgrund der hohen gesetzlichen Sicherheitsanforderungen und vor allem wegen der in der Bevölkerung vorherrschenden mangelnden Akzeptanz der Grünen Gentechnik unmöglich geworden. Wir gehen davon aus, dass sich unsere gentechnisch-ökologischen Experimente daher bis auf weiteres auf die USA beschränken werden. Nicht zuletzt auch deswegen, weil die dortige Genehmigungsbehörde, der *Animal and Plant Health Inspection Service*, Erich Kästner meistens korrekt zitiert: »Es gibt nichts Gutes, außer man tut es«.



Abbildung 11: Einfache Maßnahme zur Verhinderung der Verbreitung transgenen Pollens des Schwarzen Nachtschattens. Oben: Pflanze mit jungen Blütenknospen. Unten: Pflanze nach Entfernung der Blütenknospen (Foto: Jan-W. Kellmann).

Literatur

- [1] Pressemeldung MPI für chemische Ökologie vom 27.8.2008; Kessler D., Gase K., Baldwin I.T. (2008) Field experiments with transformed plants reveal the sense of floral scents. *Science* 321, S. 1200-1202.
- [2] Baldwin I.T., Staszak-Kozinski L., Davidson R. (1994) Up in smoke: I. smoke-derived germination cues for postfire annual, *Nicotiana attenuata* TORR. EX. WATSON; *Journal of Chemical Ecology* 20(9), S. 2345-2371.
- [3] Gaquerel E., Weinhold A., Baldwin I.T. (2009) Molecular interactions between the specialist herbivore *Manduca sexta* (Lepidoptera, Sphingidae) and its natural host *Nicotiana attenuata*. VIII. An unbiased GCxGC-ToF-MS analysis of the plant's elicited volatile emissions. *Plant Physiology* 149(3), S. 1408-1423.
- [4] Allmann S., Baldwin I.T. (2010) Insects betray themselves in nature to predators by rapid isomerization of green leaf volatiles. *Science* 329, S. 1075-1078.
- [5] Halitschke R., Baldwin I.T. (2003) Antisense LOX expression increases herbivore performance by decreasing defense responses and inhibiting growth-related transcriptional reorganization in *Nicotiana attenuata*. *The Plant Journal* 36(6), S. 794-807.
- [6] Kessler A., Halitschke R., Baldwin I.T. (2004) Silencing the jasmonate cascade: Induced plant defenses and insect populations. *Science* 305(5684), S. 665-668.
- [7] Dieser Satz und der folgende Absatz sind entnommen aus der Pressemeldung des MPI für chemische Ökologie vom 2. Juli 2004.
- [8] Hartl M., Giri A.P., Kaur H. Baldwin I.T. (2010) Serine protease inhibitors specifically defend *Solanum nigrum* against generalist herbivores but do not influence plant growth and development. *The Plant Cell* 22, S. 4158-4175.

Christoph Marksches

Wilhelm und Alexander

Das Verhältnis Alexanders zu seinem älteren Bruder Wilhelm von Humboldt¹

Über das Verhältnis der beiden Humboldt-Brüder, die nach vierzehn Jahren Debatten, ja Querelen auf Initiative liberaler Bürger seit 1883 vor dem Hauptgebäude der Berliner Humboldt-Universität und damit direkt gegenüber der damaligen Kaiserwohnung standen (und bis heute stehen),² ist überraschend wenig geschrieben worden, da und dort ein paar Zeilen und einige Fußnoten.³ Umso mehr lohnt eine Beschäftigung mit dem Thema, so weit das die – um dies gleich einschränkend anzumerken – nicht gerade abundante Quellenlage überhaupt zulässt: 1880 erschien der Briefwechsel Alexanders mit Wilhelm,

- 1 C. Marksches, »Wer nichts, als ewig todes Wissen treibet...« Über das zwiespältige Verhältnis der Brüder Humboldt. Zeitschrift für Ideengeschichte Heft IV / Frühjahr 2010, S. 39-46. Abdruck mit freundlicher Genehmigung.
- 2 K.-D. Gandert, Vom Prinzenpalais zur Humboldt-Universität. Die historische Entwicklung des Universitätsgebäudes in Berlin mit seinen Gartenanlagen und Denkmälern, Berlin 2004, 171-174. – Ich danke meiner Mitarbeiterin Marietheres Döhler und den Kolleginnen und Kollegen der Alexander-von-Humboldt-Forschungsstelle und der Wilhelm-von-Humboldt-Forschungsstelle der Berlin-Brandenburgischen Akademie der Wissenschaften, daß sie meinen eher laienhaften Kenntnissen durch Rat und Tat aufgeholfen und es mir so ermöglicht haben, den vorliegenden Text auf dem Festsymposium des Ordens Pour le mérite zum Gedächtnis an seinen ersten Ordenskanzler Alexander von Humboldt am 7. Juni 2009 vorzutragen. Dem seinerzeitigen Ordenskanzler, Herrn Kollegen Albach, danke ich für die Einladung, seinem Nachfolger Eberhard Jüngel ebenso herzlich für aufmunternde Worte.
- 3 Vgl. beispielsweise M. Geier, Die Brüder Humboldt. Eine Biographie, Reinbek bei Hamburg 2009, 270 f. oder A. Harnack, Geschichte der Königlich Preußischen Akademie der Wissenschaften zu Berlin, Bd. 1/2 Vom Tode Friedrichs des Großen bis zur Gegenwart, Hildesheim / New York 1970 (= Berlin 1900), 554 f., 571-576, 836-843 und die Literaturhinweise der folgenden Anmerkung.

und bereits im Vorwort bedauert die herausgebende Familie, daß die Gegenbriefe Wilhelms nicht erhalten seien und dazu nur ein geringer Teil der Briefe Alexanders – bekanntlich vernichtete Alexander Briefe, sofern sie für ihn nicht wissenschaftlich zu verwerten waren oder irgendwo eingeklebt werden konnten.⁴ Daher muß man auch andere Texte heranziehen, vor allem den Briefwechsel zwischen Wilhelm und Caroline, aber auch Alexanders Vorworte zur ersten Gesamtausgabe seines Bruders, die er im Übrigen intensiv begleitete,⁵ oder zu den Sonetten Wilhelms, und berücksichtigt, daß er schon 1807 »seinem teuren Bruder ... in Rom« die »Ansichten der Natur« widmete,⁶ während Wilhelm 1799 während der Lateinamerikareise Alexanders

- 4 Briefe Alexander von Humboldts an seinen Bruder Wilhelm, hg. v. d. Familie von Humboldt in Ottmachau, Stuttgart 1880, IV. Ulrich von Heinz (Berlin-Tegel) informiert mich (brieflich, 5. Juni 2009), daß die Ausgabe der Briefe in Ottmachau (Otmuchów), dem Dotationsgut Wilhelms, wo sich die Texte damals befanden, vermutlich durch den Sohn Hermann von Humboldt-Dacheröden (1809-1870) unter Mitwirkung der Enkelin Mathilde (aaO. LXXXIV) zusammengestellt wurde. Eine zweite Ausgabe (Berlin 1923) enthält leider wenig befriedigende Übersetzungen der französischen Originale. Über die hier publizierten Briefe hinaus tauchen gelegentlich Autographen im Handel auf. Vgl. auch U. v. Heinz, Die Brüder Wilhelm und Alexander von Humboldt, in: Alexander von Humboldt in Berlin. Sein Einfluß auf die Entwicklung der Wissenschaften. Beiträge zu einem Symposium, hg. v. J. Hamel, Algorismus 41, Augsburg 2003, 15-26; S.A. Kaehler, Die Brüder von Humboldt in den Jahren der napoleonischen Krise, HZ 116, 3. Folge, 20. Bd., (1916), 231-270 = ders., Studien zur deutschen Geschichte des 19. und 20. Jahrhunderts. Aufsätze und Vorträge, Göttingen 1960, 9-33; R.R. Wuthenow, Wilhelm und Alexander von Humboldt, in: Deutsche Brüder. Zwölf Doppelportraits, hg. v. Th. Karlauf u. F. Seibt, Reinbek 1994, 129-163; sowie R. Vierhaus, Die Brüder Humboldt, in: Deutsche Erinnerungsorte, Bd. 3, hg. v. E. François u. H. Schulze, München 2001, 9-25, 689 f.
- 5 A.v. Humboldt, Vorwort, in: Wilhelm von Humboldts gesammelte Werke, 1. Bd., Berlin 1841, III-VI. Humboldt schreibt, daß er »den sehnlichsten Wunsch« hatte, »diese Aufsätze bei dem Leben des Verfassers und unter seiner Mitwirkung zu sammeln; aber ein nicht zu unterdrückendes Streben nach Gediegenheit und Vollendung, wie die Strenge, mit der hochbegabte Geister ihre eigenen Schöpfungen beurtheilen, vereitelten diese Hoffnung« (aaO. III). Die vorgelegten Texte des Bruders »offenbaren uns den Menschen in dem ganzen Reichthum seines herrlichen Gemüthes und seiner Seelenkraft« (IV).
- 6 A.v. Humboldt, Ansichten der Natur, erster und zweiter Band, hg. u. kommentiert v. H. Beck, Darmstädter Ausgabe V, Darmstadt ²2008.

gesammelte »Versuche über die chemische Zerlegung des Luftkreises« herausgab.⁷ Hier kommt es, wenn wir vor allem auf der Basis von Vorworten und Briefen nach dem Verhältnis der beiden Brüder fragen, natürlich weniger auf den Klatsch an, sondern mehr auf das Verhältnis zweier wissenschaftlicher Biographien und des jeweiligen Zugriffs auf die Welt, der sich – um die These gleich vorwegzunehmen – durch eine unterschiedliche, aber dann doch überraschend kongruente Verliebtheit in den Gedanken der Totalität auszeichnet.⁸

Die beiden Brüder lagen hinsichtlich ihres Alters bekanntlich nicht sonderlich weit auseinander, wenig mehr als zwei Jahre trennten sie, allerdings überlebte Alexander seinen älteren Bruder um vierundzwanzig Jahre. »Genialer und empfindsamer, sinnlicher und schneller Dinge und Menschen erfassend erschien anfangs der ältere, langsam, kränzlich, minder erregbar der jüngere Bruder, doch selbstgefälliger und ehrgeiziger«, charakterisiert die Familie rund zwanzig Jahre nach dem Tode Alexanders das Paar.⁹ Dem entsprechen übrigens interessanterweise noch die Bewertungen der Brüder in biographischer Literatur des neunzehnten Jahrhunderts, beispielsweise bei dem Publizisten Gustav Schlesier (1810 bis nach 1853), der Alexander (im Unterschied zu heutigen Wertungen) als ein wenig geltungssüchtig und darin oberflächlich, Wilhelm aber im Kontrast als tiefsinnig porträtiert.¹⁰ Mit aller Vorsicht kann man sagen, daß Alexander die Zusammenhänge offenbar nicht gänzlich anders beurteilte: Im bereits erwähnten Vorwort zur Gesamtausgabe der Werke seines Bruders Wilhelm schreibt Alexander 1841, daß dessen »eigenthümliche Grösse ... nicht

7 A.v. Humboldt, Versuche über die chemische Zerlegung des Luftkreises und über einige andere Gegenstände der Naturlehre, Braunschweig 1799; zur Rolle Humboldts v. Heinz, Die Brüder Wilhelm und Alexander von Humboldt, 17.

8 An dieser Stelle wird vielleicht auch deutlich, warum sich der Autor dieser Zeilen, ein evangelischer Theologe, der sich normalerweise mit der Erforschung des antiken Christentums beschäftigt, für die beiden Brüder und ihr Verhältnis interessiert.

9 Briefe Alexander von Humboldts an seinen Bruder Wilhelm, XII.

10 G. Schlesier, Erinnerungen an Wilhelm von Humboldt. Neue Ausgabe, Stuttgart 1854; zu Alexander vgl. I. Schwarz, »Es ist meine Art, einen und denselben Gegenstand zu verfolgen, bis ich ihn aufgeklärt habe«. Äußerungen Alexander von Humboldts über sich selbst, in: Humboldt im Netz (HiN) 1, 2000, zitiert nach der URL: <http://www.uni-potsdam.de/u/romanistik/humboldt/hin/schwarz.htm>.

aus intellektuellen Anlagen allein, sondern vorzugsweise aus der Grösse des Charakters, aus einem von der Gegenwart nie beschränkten Sinne und aus den unergründeten Tiefen der Gefühle entspringt«. ¹¹ Ein »edles, still bewegtes Seelenleben« bescheinigt Alexander dem Bruder 1853 im gleichfalls erwähnten Vorwort seiner Edition der zu Lebzeiten nicht publizierten Sonette Wilhelms, dazu Milde und inneren Frieden des Gemüts, als er »nach langer und angestrengter Thätigkeit in einen engen Familienkreis zurücktritt, um dem Genuß der freien Natur, um großen, aber schmerzlichen Erinnerungen, um dem Studium des Alterthums und der Sprachorganismen zu leben«. ¹² Es überrascht, daß Alexander in seinem Vorwort zu den Sonetten ohne Umschweife freilich auch festhält, daß es dem Bruder in diesen Texten aus den Jahren 1831 bis 1835 an »Muße und augenblicklich auch an Neigung fehlte, in das tiefe Geheimnis von dem Verhältniß des Rhythmus zu den Gedanken einzudringen: so mußte allerdings eine mindere Sorgfalt, auf die Form gewandt, Störung des Eindrucks da verursachen, wo sich der dichterische Stoff in allzu reicher Fülle dargeboten hatte«. ¹³ Abgemildert wird die kritische Note dann durch Hinweise auf die metrische Übersetzung von Aischylos' Agamemnon von 1816 und andere Arbeiten Wilhelms zum Versmaß der griechischen Tragödie. ¹⁴ 1837 kritisierte Alexander in einem Brief an Varnhagen recht deutlich den 1821 in der Akademie vorgetragenen Essay seines Bruders über »Die Aufgabe des Geschichtsschreibers«. ¹⁵ Wilhelms Versuch, dem Geschichtsschreiber neben der Rekonstruktion des Geschehenen die Darstellung der »Ideen als ›Gesetzen der notwendigen Kette des Wirklichen‹« zuzuweisen und zu argumentieren, daß die Weltgeschichte »nicht ohne eine Weltregierung verständlich« sei, ¹⁶ anerkennt Alexander (»Meines Bruders Aufsatz gehört zu dem Vollendetsten in Sprache, das er ge-

11 A.v. Humboldt, Vorwort, in: Wilhelm von Humboldts gesammelte Werke, 1. Bd., Berlin 1841, IV.

12 A.v. Humboldt, Vorwort, in: Sonette von Wilhelm von Humboldt, Berlin 1853, (III-XVI) IIIf.

13 A.v. Humboldt, Vorwort, in: Sonette, VI f.

14 Aeschylos Agamemnon (1816), in: Wilhelm von Humboldts gesammelte Werke, 3. Bd. Berlin 1843, 1-96.

15 W.v. Humboldt, Über die Aufgabe des Geschichtsschreibers, in: Wilhelm von Humboldt, Werke in fünf Bänden, hg. v. A. Flitner u. K. Giel, Bd. 1 Schriften zur Anthropologie und Geschichte, Darmstadt 3 1980, 585-606.

16 W.v. Humboldt, Über die Aufgabe des Geschichtsschreibers, 587, 600.

schrieben«). Aber er fügt die Kritik hinzu, indem er die Grundidee des Essays, wie mir jedenfalls scheint, mehr als erlaubt zuspitzt: »Gott regiert die Welt; die Geschichtsaufgabe ist das Aufspüren dieser ewigen geheimnisvollen Rathschlüsse, das ist doch eigentlich das Resultat, und über dies Resultat habe ich bisweilen mit meinem Bruder, ich darf nicht sagen gehadert, sondern diskutiert. ... Ich weiß, Sie werden mir zürnen, weil Sie erraten, daß die Hauptidee dieser herrlichen Abhandlung mich nicht ganz befriedigt.«¹⁷ Umgekehrt waren Wilhelm und Caroline von Dacheröden vor allem zu Beginn ihrer Beziehung stellenweise recht kritisch mit dem Bruder, jedenfalls ziemlich schwankend im Urteil; Wilhelm schreibt beispielsweise aus Berlin 1790 an Georg Forster über Alexander: »Suchen Sie ihm, wenn Sie Gelegenheit haben, etwas mehr Selbstzufriedenheit zu geben. Er hat ihrer wirklich von gewissen Seiten zu wenig und das macht ihn unglücklich.«¹⁸ Und Caroline schließt ihre Beschreibung der Kosmos-Vorlesungen in der Singakademie für die Tochter Adelheid im Dezember 1827 mit dem Ausruf über den Schwager: »Ach, und doch nicht glücklich!«.¹⁹ Man stand sich also bei aller persönlichen Nähe durchaus nicht kritiklos gegenüber. Wilhelm schrieb einmal über den Bruder: »Man kommt der Natur, sieht man daraus, nicht näher, wenn man aus der zivilisierten Welt hinausgeht«, ²⁰ und reagierte stellenweise äußerst befremdet auf das, was er und Caroline als französische Allüren des aus Amerika heimgekehrten weltberühmten Forschungsreisenden deuteten: »Vor der Welt«, so schreibt Wilhelm an Caroline in Paris im August 1804, »muß man das Vaterland ehren, wenn es auch eine Sandwüste ist.«²¹

17 Brief A.v. Humboldts an Varnhagen vom 10.5.1837, in: Briefe von Alexander von Humboldt an Varnhagen von Ense aus den Jahren 1827 bis 1858, nebst Auszügen aus Varnhagens Tagebüchern, und Briefen von Varnhagen und Andern an Humboldt (hg. v. L. Assing), Leipzig 1860, 40.

18 Brief W.v. Humboldts an Georg Forster vom 5.4.1790, in: Briefe an Forster, Georg Forsters Werke XVIII, Berlin 1982, nr. 265, p. 389 f.

19 Brief Carolines an Adelheid von Humboldt-Dacheröden vom 7.12.1827, zitiert nach: Wilhelm und Caroline von Humboldt in ihren Briefen, 7. Bd. Reife Seelen. Briefe von 1820-1835, Berlin 1916, 325.

20 Brief W.v. Humboldts an Caroline vom 6.6.1804, in: Wilhelm und Caroline von Humboldt in ihren Briefen, hg. v. A.v. Sydow, 2. Bd. Von der Vermählung bis zu Humboldts Scheiden aus Rom 1791-1808, Berlin 1907, 183.

21 Brief W.v. Humboldts an Caroline vom 29.8.1804, in: Wilhelm und Caroline von Humboldt in ihren Briefen, hg. v. A.v. Sydow, 2. Bd., 232; vgl. auch M. Geier, Die Brüder Humboldt, 241-251.

Spätestens im Oktober 1813 fanden es die beiden nur noch höchst unpassend, daß Alexander in Zeiten wie diesen in Paris lebte²² – Alexander selbst bezeichnete sich bekanntlich einmal in einem Brief an Christian Carl Josias von Bunsen aus dem Jahr 1842 als »ein alter tri-colorer Lappen ..., den man conservirt, und der (kommt einmal die Noth wieder) deployirt [sc. entfaltet] werden kann«.²³ Wilhelm berichtete Caroline im Dezember 1813, daß der kleine Prinz Wilhelm von Preußen gefragt habe, wo Alexander von Humboldt jetzt sei, und auf die Antwort »in Paris« lebhaft geantwortet habe: »Gott! Da sollte er Napoleon totschiessen«. Und den Brief an die Gattin schließt Wilhelm von Humboldt: »Wäre das nicht ein hübscher Auftrag für Alexander?«.²⁴ Alexanders mangelndes Talent, mit Geld umzugehen, und seine teils erheblichen Schulden erregen auch immer wieder den Unwillen seines Bruders, natürlich besonders, wenn er selbst in diese Geld- und Leihgeschäfte verwickelt ist.²⁵ Und es ist immer wieder behauptet worden, daß Wilhelm von Humboldt bei einem berühmten Sonett aus dem Jahre 1834 über das »tote Wissen« vor allem an seinen Bruder Alexander dachte:

Wer nichts, als ewig todes Wissen treibet,
Niemals Gedanken zur Idee verbindet,
Nie, was den Busen tief ergreift, empfindet,
Der doch nur in der Menschheit Vorhof bleibt.

Freilich war damals der »Kosmos« noch nicht veröffentlicht, und man kann das Sonett vielleicht auch als etwas scharfe Aufforderung deuten,

22 Brief W.v. Humboldts an Caroline vom 6.12.1813, in: Wilhelm und Caroline von Humboldt in ihren Briefen, hg. v. A.v. Sydow, 4. Bd. Federn und Schwerter in den Freiheitskriegen. Briefe von 1812–1815, Berlin 1910, 188 (»Ja, ich gestehe Dir frei, was ich sonst nicht sage, daß ich auch an Alexander sein Bleiben in Paris nicht billige«) und noch schroffer in einem Brief an die Frau vom 12.11.1817 aus London, in: 6. Bd. Im Kampf mit Hardenberg. Briefe von 1817–1819, Berlin 1913, 46 f.; vgl. H. Rosenstrauch, Wahlverwandt und ebenbürtig. Caroline und Wilhelm von Humboldt, Die andere Bibliothek, Frankfurt/Main 2009, 215.

23 Brief A.v. Humboldts an Bunsen, in: Briefe von Alexander von Humboldt an Christian Carl Josias Freiherr von Bunsen, Leipzig 1869, 51.

24 Brief W.v. Humboldts an Caroline vom 9.12.1813, in: Wilhelm und Caroline von Humboldt in ihren Briefen, hg. v. A.v. Sydow, 4. Bd., 192.

25 Vgl. z.B. W.v. Humboldt an Caroline am 2.5.1809, in: Wilhelm und Caroline von Humboldt in ihren Briefen, hg. v. A.v. Sydow, 3. Bd., 152.

das Projekt endlich abzuschließen.²⁶ Schön ist auch folgende Charakterisierung Wilhelms aus dem Jahre 1817: »Wir sind uns sehr gut, aber selten einig. Darum sprechen wir auch wenig zusammen. Unser Charakter ist zu verschieden.«²⁷

Die ausführliche Biographie beider Brüder, die zeitweilig äußerst enge Verbindung ihrer beiden Lebensläufe muß uns nicht länger beschäftigen, sie ist bekannt und angesichts des täglichen Umgangs auch nicht immer besonders leicht im Detail zu beschreiben: 1786/1787 gemeinsam bei Henriette Herz, 1789/1790 Zeiten gemeinsamen Studiums in Göttingen, 1797 gemeinsame Aufenthalte in Jena und Weimar, nach der großen Reise ein Treffen in Rom im Frühsommer 1805, gelegentliche weitere Treffen und Besuche bis zur endgültigen Übersiedlung nach Berlin im Mai 1827 (»wo ich endlich das lang entbehrte Glück genoß, mit meinem Bruder an einem Orte zu leben und vereint wissenschaftlich zu arbeiten«, wie Alexander 1850 schrieb),²⁸ die Zäsur durch Carolines Tod im März 1829²⁹ und schlußendlich Wilhelms Tod im April 1835. In einem Brief an Varnhagen in jenem April 1835, drei

26 Sonett Nr. 1043, in: Wilhelm von Humboldts Gesammelte Schriften, Bd. IX, Berlin 1903, 421 f.; vgl. allerdings v. Heinz, Die Brüder Wilhelm und Alexander von Humboldt, 19, kritischer S.A. Kaehler, Wilhelm von Humboldt und der Staat. Ein Beitrag zur Geschichte deutscher Lebensgestaltung um 1800, Göttingen 1963, 40 f., der die Beziehung des Sonetts auf Humboldt mit seinem Herausgeber, dem Berliner Germanisten Albert Leitzmann (1867-1950), für »einwandfrei« hält (aaO. Anm. 1 zu S. 40 auf S. 447).

27 Zitiert nach: R.R. Wuthenow, Wilhelm und Alexander von Humboldt, 141.

28 In einem Text für den Verleger Brockhaus: Aus meinem Leben (1769-1850), in: Alexander von Humboldt, Aus meinem Leben. Autobiographische Bekenntnisse, zusammengestellt und erläutert v. K.-R. Biermann, München 1989, (83-119) 115 f. – Schon 1790 kann man sich vorstellen, zusammenzuleben: »Wenn er einmal mit uns leben könnte, sollt es mich recht freuen« (Caroline an Wilhelm von Humboldt am 1.5.1790, in: Wilhelm und Caroline von Humboldt in ihren Briefen, hg. v. A.v. Sydow, 1. Bd. Briefe aus der Brautzeit 1787-1791, Berlin 1910, 143).

29 Damals schreibt Alexander Wilhelm: »Niemand in dieser Welt liebt Dich so zärtlich wie ich. Meine Existenz wird für immer an Deine geknüpft sein; und wir wollen uns niemals mehr auf lange trennen«, (St. Petersburg, 9./21.11.1829), hier zitiert nach: Wuthenow, Wilhelm und Alexander von Humboldt, 160.

Tage vor dem Tod des Bruders, beschreibt Alexander die letzten Stunden des Parkinsonkranken: »Immer noch die Klarheit des großen Geistes« und zitiert letzte Worte: »Ich war sehr glücklich: ... denn die Liebe ist das Höchste. Bald werde ich bei der Mutter sein, Einsicht haben in eine höhere Weltordnung«.³⁰

Für die vorausgehenden Jahre gilt: Sah man sich nicht, vermißte man sich, teils heftig: »Es tut mir sehr leid, mich so lange von ihm getrennt zu sehen, allein es ist eine schöne Reise, er ist ganz dazu gemacht, sie, so wie es geschehen muß, zu benutzen, und so teile ich seine in der Tat außerordentliche Freude«, schrieb Wilhelm, nachdem der Bruder nach Lateinamerika aufgebrochen war, im April 1799 an Schiller.³¹ Und in Alexanders 1801 in Santa Fé de Bogotá/Kolumbien entstandener, nicht zur Veröffentlichung bestimmter Lebensbeschreibung findet sich über die Jahre, die der großen Lateinamerika-Reise vorangingen, folgende Formulierung des jüngeren der beiden Brüder: »Alles, was auf bürgerliche Verhältnisse Bezug hatte, wurde mir verächtlich, jede Gemächlichkeit des häuslichen Lebens und der feineren Welt ekelte mich an. ... Wilhelms Abwesenheit (er war in Paris mit Campe) vermehrte die Krisis«.³² Entsprechend waren dann die Zusammentreffen der beiden Brüder höchst intensiv: In »Zerstreuung und Geisteszerrüttung« versetzt Wilhelm im März 1796 die Anwesenheit seines Bruders Alexander in Tegel,³³ seine Anwesenheit stört »gleich sehr in meiner Assiette und meinen Studien«, aber sie ist »doch äußerst fruchtbar an mir noch fremden oder nur halb bekannten Ideen gewesen«. Wenn man sich nicht sah, schrieb man Briefe. Insbesondere die langen, teils schon zu Lebzeiten in Berlin veröffentlichten Texte von der Reise nach Lateinamerika enthalten beeindruckende Schilderungen von Natur wie Kultur und mühsam überlebten Gefahren »unter den Indianern sowohl als in den mit Krokodilen, Schlangen und

30 Brief A.v. Humboldts an Varnhagen vom 5.4.1835, in: Briefe von Alexander von Humboldt an Varnhagen von Ense aus den Jahren 1827 bis 1858, 26.

31 Brief an Schiller vom 26.4.1799, zitiert nach: Der Briefwechsel zwischen Friedrich Schiller und Wilhelm von Humboldt, Bd. 2, Berlin 1962, 178.

32 A.v. Humboldt, Ich über mich selbst, in: Alexander von Humboldt, Aus meinem Leben. Autobiographische Bekenntnisse, (31-40) 40.

33 Brief an Schiller vom 2.3.1796, zitiert nach: Der Briefwechsel zwischen Friedrich Schiller und Wilhelm von Humboldt, Bd. 2, 41 f.

Tigern angefüllten Wüsten«. ³⁴ Die von der Familie bewahrten Briefe an Wilhelm stellen nur einen Ausschnitt dar. An einem bislang nicht veröffentlichten Brief Alexanders im Archiv der Berlin-Brandenburgischen Akademie der Wissenschaften, der vermutlich aus dem Jahre 1832 stammt, wird deutlich, daß man auch über alltägliche Probleme der wissenschaftlichen Arbeit korrespondierte. So beschäftigt sich Alexander, der wahrscheinlich gerade wieder in Paris weilt, in diesem mehr zufällig erhaltenen Brief mit zwei Sanskrit-Worten, die der antike Geograph Ptolemaeus überliefert, und weiteren Begriffen dieser Sprache, die sich in der Kommentierung des spätantiken christlichen Weltreisenden Kosmas Indicopleustes finden, nimmt aber alle Gelehrsamkeit bei deren Übersetzung sofort ironisch zurück: »Du siehst, ich kann sehr langweilig werden«, und fügt den offenbar weniger langweiligen Berliner Hofkatsch an: »Gewisser ist, daß die Königin der Niederlande Krämpfe von einer falsch verschluckten Auster hatte. Auch die Meerthiere empören sich gegen die Meeres Königin«.

Anläßlich der Einweihung der eingangs erwähnten Denkmäler der Brüder Humboldt vor der Berliner Universität sprach in deren Hauptgebäude 1883 Emil du Bois-Reymond, einst Korrespondenzpartner Alexanders. Nach einem kurzen Durchgang durch die Geschichte der Errichtung der Denkmäler wandte er sich den beiden Dargestellten zu und bemerkte eingangs, daß selbst »auf einem und demselben Gebiet vergleichende Schätzungen geistiger Größe äußerst mißlich« seien, »vollends wenn es um verschiedenartige Gaben und Leistungen sich handelt, für die es kein gemeinsames Maß gibt«. ³⁵ Und du Bois-Reymond sagte weiter: »Doch wir wollen den unfruchtbaren Streit, wer von den Brüdern eigentlich der größere war, nicht aufnehmen, sondern

34 Brief Alexanders an Wilhelm vom 17.10.1800, in: Briefe Alexander von Humboldts an seinen Bruder Wilhelm, 17.

35 E. du Bois-Reymond, Die Humboldt-Denkmäler vor der Berliner Universität, in der Aula der Berliner Universität am 3. August 1883 gehaltene Rede, in: Briefwechsel zwischen Alexander von Humboldt und Emil du Bois-Reymond, hg. v. I. Schwarz u. K. Wenig, Beiträge zur Alexander-von-Humboldt-Forschung 22, Berlin 1997, (185-203) 188. Vgl. aber ebd.: »Wilhelms Nachruhm kam es zugute, daß er, auf der Höhe hinweggerafft, schon heroisiert wurde zu einer Zeit, wo Alexander noch als erhabene Ruine unter den Lebenden weilte, und durch manche greisenhafte Schwäche seinem Ansehen in der Nähe schadete«.

uns des günstigen Geschicks freuen, welches sie hier vereint. Nur eines dürfte sicher für Alexander zu beanspruchen sein: daß er, um zu werden, was er der Welt war, kühner voranzugehen und einen längeren und schwierigeren Weg zurückzulegen hatte als, um zu seinem Ziele zu gelangen, sein Bruder Wilhelm«. ³⁶ Mir scheint, daß dieser Satz nicht nur zutrifft, wenn man ihn im schlichtesten Sinne seines Wortlauts als Beschreibung der langen Reisen Alexander von Humboldts interpretiert; Wilhelm hat das in einem Brief an Caroline bereits rund achtzig Jahre vor du Bois-Reymond ähnlich formuliert. ³⁷

Nun war eingangs behauptet worden, daß die beiden ungeachtet aller charakterlichen und in ihren Arbeitsfeldern liegenden Unterschiede so ähnlichen Brüder durch eine (in Details ebenfalls unterschiedliche, aber dann doch überraschend kongruente) Verliebtheit in den Gedanken der Totalität geprägt sind – oder, um präziser zu reden: in den Gedanken der Totalität, den sie mit dem Gedanken der Individualität in Balance zu bringen versuchen. In den Prolegomena zu Alexanders »Kosmos« heißt es beispielsweise: »Was mir den Hauptantrieb gewährte, war das Bestreben, die Erscheinungen der körperlichen Dinge in ihrem allgemeinen Zusammenhang, die Natur als ein durch innere Kräfte bewegtes und belebtes Ganzes aufzufassen«, ³⁸ entsprechend lautet das aus der Naturgeschichte des Plinius entnommene Motto des Gesamtwerks, das Eberhard Knobloch in einem Aufsatz, in dem weitere Linien zwischen Plinius und Humboldt gezogen sind, glücklich wie folgt ins Deutsche übersetzt hat: »Aber die Kraft und die Großartigkeit der Dinge der Natur entbehren in all ihren Wechseln der Glaubwürdigkeit, wenn jemand im Geiste nur deren Teile und sie

36 Du Bois-Reymond, Die Humboldt-Denkmäler vor der Berliner Universität, 189.

37 Brief W.v. Humboldts an Caroline vom 9.10.1804, in: Wilhelm und Caroline von Humboldt in ihren Briefen, 2. Bd., 260: »Ich kann nicht sagen, daß ich eben glaube, in irgend einer Art über ihm zu stehen. Aber seit unserer Kindheit sind wir wie zwei entgegengesetzte Pole auseinandergegangen, obgleich wir uns immer geliebt haben und sogar vertraut miteinander gewesen sind. Er hat von früh an nach außen gestrebt, und ich habe mir ganz früh schon nur ein inneres Leben erwählt, und glaube mir, darin liegt alles«.

38 A.v. Humboldt, Kosmos. Entwurf einer physischen Weltbeschreibung, hg. u. kommentiert v. H. Beck, Teilband 1, Alexander-von-Humboldt-Studienausgabe VII/1, Darmstadt 1993 (= Stuttgart und Augsburg 1845), 7.

nicht als ganze erfaßt«. ³⁹ Und Humboldt formuliert weiter: »Ein beschreibendes Naturgemälde, wie wir es in diesen Prolegomenen aufstellen, soll aber nicht bloß dem Einzelnen nachspüren; es bedarf nicht zu seiner Vollständigkeit der Aufzählung aller Lebensgestalten, aller Naturdinge und Naturprozesse. Der Tendenz endloser Zersplitterung des Erkannten und Gesammelten widerstrebend, soll der ordnende Denker trachten, der Gefahr der empirischen Fülle zu entgehen«. ⁴⁰ Ob man im Blick auf diesen »Ganzheitsgedanken« tatsächlich – wie Hartmut Böhme meint – von einer »eigenartigen Entwicklungslosigkeit« im Œuvre Alexander von Humboldts sprechen kann, ⁴¹ können wir hier dahingestellt sein lassen, ich glaube das nicht. ⁴² Außerdem ist sicher, daß Wilhelm sich seinen Bruder so wünschte, wie er uns in den ersten Kosmos-Bänden begegnet. Schon 1793 schrieb er über Alexander: »Er ist gemacht, Ideen zu verbinden, Ketten von Dingen zu erblicken, die Menschenalter hindurch ohne ihn unentdeckt geblieben wären«. Und dazu ist es notwendig, »Einheit in alles menschliche Streben zu bringen, zu zeigen, daß diese Einheit der Mensch ist, und zwar der innere Mensch ist, und den Menschen zu schildern, wie er auf alles außer ihm, und wie alles außer ihm auf ihn wirkt, daraus

39 Plin., hist. nat. VII 1: *Naturae vero rerum vis atque maiestas in omnibus momentis fide caret, si quis modo partes eius ac non totam complectatur animo* (zitiert nach Kosmos, Studienausgabe VII/1, 5; zur Übersetzung vgl. E. Knobloch, Naturgenuss und Weltgemälde. Gedanken zu Humboldts Kosmos, in: Dahlemer Archivgespräche 12, 2006, [24-41] 28). In einem Brief an Varnhagen vom 27.10.1834 bemerkt Humboldt: »Ich habe den tollen Einfall, die ganze materielle Welt, alles was wir heute von den Erscheinungen der Himmelsräume und des Erdenlebens, von den Nebelsternen bis zur Geographie der Moose auf den Granitfelsen wissen, alles in Einem Werke darzustellen, und in einem Werke, das zugleich in lebendiger Sprache anregt und das Gemüth ergötzt. Jede große und wichtige Idee, die irgendwo aufglimmt, muß neben den Thatsachen hier verzeichnet sein« (Briefe von Alexander von Humboldt an Varnhagen von Ense aus den Jahren 1827 bis 1858, 20). Laut dem Brief präferierte auch Wilhelm von Humboldt den Titel »Kosmos« für das Werk (aaO. 22).

40 Humboldt, Kosmos, Studienausgabe VII/1, 62.

41 H. Böhme, Ästhetische Wissenschaft. Aporien im Werk Alexander von Humboldts, in: Alexander von Humboldt – Aufbruch in die Moderne, hg. v. O. Ette u.a., Beiträge zur Alexander-von-Humboldt-Forschung 21, Berlin 2001, 17-33.

42 Vgl. dazu unten Anm. 46.

den Zustand des Menschengeschlechts zu zeichnen ... Das Studium der physischen Natur nun mit dem der moralischen zu verknüpfen, und in das Universum, wie wir es kennen, eigentlich erst die wahre Harmonie zu bringen, oder wenn dieß die Kräfte Eines Menschen übersteigen sollte, das Studium der physischen Natur so vorzubereiten, daß dieser zweite Schritt leicht werde, dazu, sage ich, hat mir unter allen Köpfen, die ich historisch und aus eigener Erfahrung in allen Zeiten kenne, nur mein Bruder fähig geschienen«. ⁴³ Bekanntlich prägen solche Totalitätsideen aber auch Wilhelm von Humboldts eigene Konzeptionen, beispielsweise im Blick auf seine Konzeption universaler Bildung, die die Totalität geistigen Vermögens zur Reife bringen soll, ⁴⁴ natürlich auch im Blick auf Kunst- und Literaturtheorie ⁴⁵ oder die Sprachphilosophie ⁴⁶ – das ist bekannt und muß nicht vertieft werden.

Kritik an solchen Totalitätsverliebtheiten ist leicht – heutigentags noch leichter als damals, und ein Theologe wird immer vor irdischen Totalitätsbehauptungen warnen. In der berühmten Konstanzer Denkschrift »Geisteswissenschaften heute« aus dem Jahr 1990 kann man lesen, daß es »angesichts über 4000 wissenschaftlicher Fächer, die der Fächerkatalog des Hochschulverbandes aufzählt«, schwer falle, »von einem System der Wissenschaft oder auch nur von einer disziplinären

43 Brief W.v. Humboldts an Brinckmann vom 18.3.1793 (Wilhelm von Humboldts Briefe an Karl Gustav von Brinckmann, hg. u. erläutert v. A. Leitzmann, Leipzig 1939, nr. 27, (p. 57-62) 61 (auch bei Kaehler, Wilhelm von Humboldt und der Staat, 41 und v. Heinz, Die Brüder Wilhelm und Alexander von Humboldt, 20 zitiert).

44 So beispielsweise im »Gutachten über die Organisation der Ober-Examinations-Kommission« von 1809, in: Wilhelm von Humboldt. Werke in fünf Bänden, Bd. IV Schriften zur Politik und zum Bildungswesen, Darmstadt ²1982, (77-89) 82-85.

45 So beispielsweise in: Aesthetische Versuche. Erster Teil. Ueber Göthes Hermann und Dorothea, in: Wilhelm von Humboldt. Werke in fünf Bänden, Bd. II Schriften zur Altertumskunde und Ästhetik. Die Vasken, Darmstadt ³1979, (125-356) 145-152.

46 W.v. Humboldt, Ueber das vergleichende Sprachstudium in Beziehung auf die verschiedenen Epochen der Sprachentwicklung, in: Wilhelm von Humboldt. Werke in fünf Bänden, Bd. III Schriften zur Sprachphilosophie, Darmstadt ³1979, (1-25) 24 (Totalität aller denkbaren Sprachen).

Ordnung der Fächer in einem solchen System zu sprechen«. ⁴⁷ Totalitätsideen treten daher gegenwärtig meist nur noch in Gestalt leicht übergreifiger Totaltheorien, beispielsweise von bestimmten Fachvertretern der Neurologie, auf. Angesichts dieser Lage wird natürlich niemand einfach für eine Erneuerung der Humboldtschen Totalitätsideen in Form bloßer Repetition plädieren (zumal sich auch bei den Brüdern an diesem Punkt Ansätze zur kritischen Reflexion finden), ⁴⁸ aber es scheint ein mindestens gelegentliches Studium dieser Ideen sinnvoll, damit die Zusammengehörigkeit aller wissenschaftlichen Erkenntnis wenigstens als regulative Idee präsent bleibt.

47 W. Frühwald, H.R. Jauß, R. Koselleck, J. Mittelstraß, B. Steinwachs, Geisteswissenschaften heute. Eine Denkschrift, Konstanz 1990 (als Manuskript gedruckt), 16.

48 Vgl. aus der Einleitung zum fünften und letzten, unvollendet gebliebenen Band des »Kosmos«: »Wenn der »Kosmos« (...) in seiner wissenschaftlichen Richtung durchaus dem Zeitalter der Empirie angehört, wenn der universalistische Charakter, den er an sich trägt, ein durchaus unspeculativer ist, vielmehr dem ästhetischen Stempel unserer literarischen Periode seinen Ursprung verdankt, so war nun nicht bloß diese literarische, sondern auch die anti-speculativ empiristische Periode bereits vorüber«, zitiert nach: Alexander von Humboldt. Eine wissenschaftliche Biographie, im Verein mit R. Avé-Lallement, J.V. Carus, A. Dove u.a. bearb. u. hg. von K. Bruhns, Leipzig 1872, 2. Bd., 416.

Jens Frahm, Martin Uecker,
Dirk Voit und Shuo Zhang

Magnetresonanz-Tomografie in Echtzeit¹

Die Magnetresonanz-Tomografie (MRT) ist ein nichtinvasives Verfahren für die diagnostische Bildgebung, das auch der breiteren Öffentlichkeit als Alternative zur Röntgen-Computertomografie (CT) bekannt ist. Die große klinische und wissenschaftliche Bedeutung der MRT gründet sich auf eine Vielzahl von unterschiedlichen Techniken, die weit über die anatomisch-strukturelle Darstellung des Gewebes – in der Regel sogar ohne Kontrastmittel – hinausgehen und beispielsweise funktionelle Untersuchungen der Durchblutung oder der Hirntätigkeit ermöglichen. Fast alle MRT-Verfahren haben jedoch einen auch für die Patienten spürbaren Nachteil: Sie sind empfindlich gegenüber Bewegungen und liefern bei entsprechenden Störungen Bilder mit fehlerhaften Darstellungen. Dynamische Messungen schneller physiologischer Vorgänge sind daher nur durch eine Synchronisierung der Datenaufnahme mit der Atmung oder der Herzbewegung möglich. Diese Probleme können nun durch Methoden überwunden werden, die unsere Arbeitsgruppe in den letzten zwei Jahren in mehreren Schritten entwickelt hat. Sie zeichnen sich durch eine Unempfindlichkeit gegenüber Bewegungen aus und ermöglichen die Aufnahme schneller Bildserien oder »MRT-Filme« mit einer bisher undenk바aren zeitlichen Auflösung von 20 Millisekunden pro Bild. Der sich daraus ergebende Zugang zu dynamischen Prozessen in Echtzeit wird die MRT und ihre Anwendungen nachhaltig verändern.

Bildgebung mit der Magnetresonanz-Tomografie

Die MRT beruht auf der Beobachtung von Radiosignalen im Frequenzbereich des Rundfunks. Sie werden von den Wasserstoffatomkernen (Protonen) des Wassers in unserem Körper ausgesendet, wenn dieser sich in einem starken Magnetfeld befindet und mit einem kur-

¹ Vortrag beim XLAB Science Festival 2010.

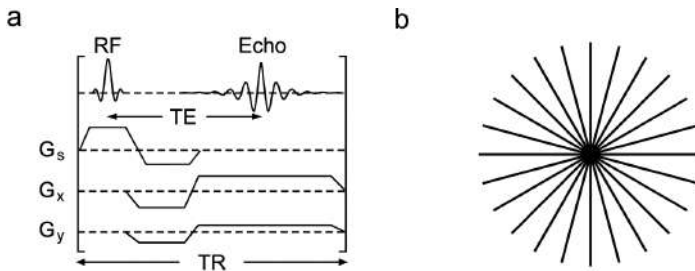


Abbildung 1: (a) FLASH MRT-Sequenz mit radialer Ortskodierung (einzelne Teilmessung): RF = Radiofrequenz-Anregungsimpuls, Echo = aufgezeichnetes Gradientenecho-Signal (Datenzeile), TE = Echozeit, G_z = schichtselektiver Magnetfeldgradient, G_x und G_y = variable frequenzkodierende Gradienten für die radiale Ortskodierung, TR = Repetitionszeit für die Wiederholung der Teilmessung mit unterschiedlicher Ortskodierung. (b) Radiale Abtastung des Datenraumes (Beispiel mit 12 Speichen).

zen Impuls aus ebensolchen Ultrakurzwellen angeregt wird. Um mit der MRT ein Bild von der Verteilung des Wassers im Gewebe zu gewinnen, werden räumlich unterscheidbare MRT-Signale durch magnetische Zusatzfelder erzeugt. Aus einer Vielzahl von Einzelmessungen mit veränderten Ortskodierungen lassen sich anschließend zwei- oder dreidimensionale Bilder des Gewebes errechnen. Diese vielen Teilmessungen sind für die relativ lange Messzeit der MRT verantwortlich. Die Kontraste in den MRT-Bildern spiegeln Unterschiede in der Konzentration und Beweglichkeit des Wassers in den verschiedenen Geweben wider, beispielsweise zwischen weißer Hirnsubstanz (Bündel von Nervenfasern) und grauer Hirnsubstanz (Hirnrinde aus den Nervenzellkörpern).

Schnelle Verfahren der MRT

Die aktuellen Entwicklungen beruhen zum Teil auf einer Schnellaufnahmetechnik, die bereits vor 25 Jahren die MRT revolutioniert hat: das FLASH-Verfahren (*Fast Low Angle SHot*) reduzierte die Messzeiten der Bilder um mindestens das Hundertfache, so dass Schnittbilder im Sekundenbereich gemessen werden konnten und erstmals dreidimensionale Aufnahmen möglich wurden (Frahm et al. 1985). Das in Abbildung 1a

schematisch dargestellte FLASH-Prinzip verwendet Radiofrequenzimpulse (RF) mit besonders kleiner Leistung, die nur einen Teil der Protonen im Gewebe anregen. Die Anwendung eines Magnetfeldgradienten (G_s) während des RF-Impulses begrenzt die Bildgebung auf eine einzelne Schicht. Zum Zeitpunkt TE nach dem RF-Impuls wird das erzeugte MRT-Signal (Echo) aufgezeichnet. Durch die gleichzeitige Schaltung von Magnetfeldgradienten (G_x und G_y) wird dieses Signal ortskodiert.

Die starke Beschleunigung der MRT-Messung geht auf die Tatsache zurück, dass bei einer Anregung des MRT-Signals mit nur einem kleinen Anregungsimpuls der größte Anteil der Protonen ungenutzt bleibt und für eine unmittelbare Fortsetzung der Messung mit dem nächsten ortskodierenden Telexperiment zur Verfügung steht. Damit kann die lange Wartezeit zwischen den Telexperimenten von mindestens einer Sekunde, die zur Erholung des MRT-Signals bei einem vollständigen Verbrauch erforderlich ist, vollständig entfallen. Die Zeit für eine Teilmessung reduzierte sich damit auf unter 10 Millisekunden.

Die jetzt weiter erhöhte Toleranz gegenüber Bewegungen beruht auf einer im Vergleich mit der üblichen MRT-Technik veränderten Art der Ortskodierung: Statt den Datenraum für ein Bild zeilenweise in einem rechtwinkligen Raster unter Nutzung von Frequenz und Phasenlage der MRT-Signale abzutasten, geschieht dies, wie in Abbildung 1b dargestellt, in radialer Form wie die Speichen eines Rades. Durch eine geeignete Kombination der beiden gleichzeitig geschalteten Magnetfeldgradienten G_x und G_y lässt sich jede beliebige Winkelposition einstellen. Für die Rekonstruktion eines MRT-Bildes aus radial kodierten Daten werden die Datenpunkte zunächst auf ein rechtwinkliges Raster interpoliert (*Gridding*) und anschließend wie bei der konventionellen MRT mittels einer zweidimensionalen Fourier-Transformation in ein Bild umgeformt.

Die hohe Unempfindlichkeit der radialen Ortskodierung gegenüber Bewegungen beruht auf der Tatsache, dass nur die Frequenzkodierung der MRT-Signale genutzt wird: Die durch Bewegungen beeinflussbare Phasenkodierung wird völlig vermieden. Abbildung 2 zeigt als Beispiel Aufnahmen des Brust- und Bauchraums, die mit der radialen FLASH-MRT bei freier Atmung und einer relativ langen Messzeit von einer Sekunde pro Bild aufgenommen wurden. Trotz der sehr viel schnelleren Bewegung des Herzens entstehen keine Fehler, die sich bei der konventionellen Technik über das gesamte Bild er-

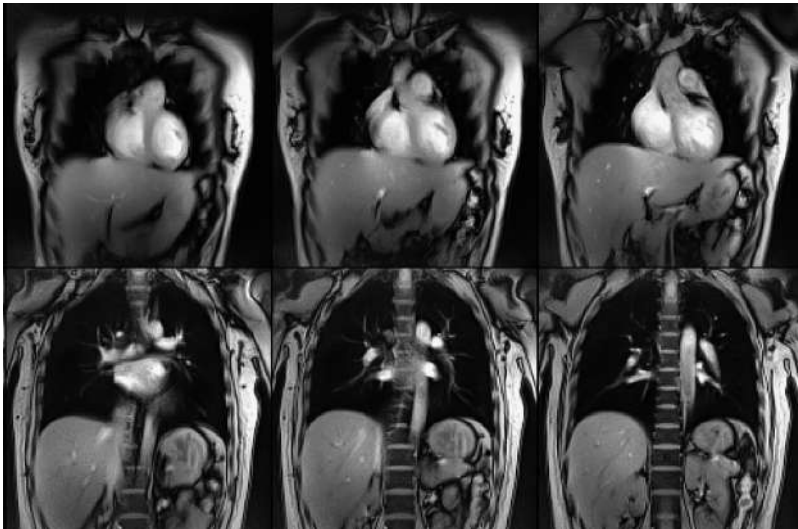


Abbildung 2: Radiale FLASH-MRT von Brust- und Bauchraum bei freier Atmung und Herzbewegung. Die einzelnen Aufnahmen mit Messzeiten von einer Sekunde zeigen zwar eine zeitliche Mittelung des sich rasch bewegenden Herzmuskels (grau) und des Blutes in den Herzkammern (hell), aber keine groben Bildfehler.

strecken würden, sondern allenfalls zeitliche Mittelwerte der Signale bewegter Objekte.

Höchstgeschwindigkeit durch Unterabtastung

Für die dynamische Bildgebung in Echtzeit ist es neben der Bewegungstoleranz wichtig, möglichst kurze Messzeiten zu erzielen. Da die physiologisch zulässigen Grenzen mit der heutigen MRT-Geräte-technik bereits erreicht werden (d.h. Teilmessungen innerhalb von einigen Millisekunden sind möglich), ist eine weitere Verkürzung der Messzeit nur noch durch eine Reduktion der für eine Bild-Rekonstruktion notwendigen Datenmenge, also eine geringere Anzahl von Teilmessungen, möglich: Je weniger Daten aufgenommen werden, desto kürzer die Messzeit. Für diese Strategie bietet die radiale Abtastung des Datenraumes einen besonderen Vorteil, der auf der Gleich-

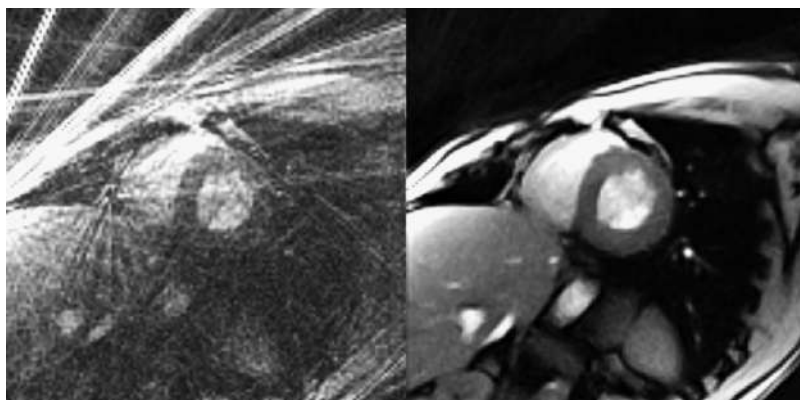


Abbildung 3: Bild-Rekonstruktion (Links) durch konventionelles *Gridding* und (Rechts) durch regularisierte nichtlineare Inversion derselben Daten einer MRT-Messung des schlagenden Herzens (postsystolische Expansionsphase, Kurzachsenblick der rechten und linken Herzkammer). Die Messzeit des Bildes betrug 30 Millisekunden bei einer Datenmenge von nur 15 Speichen.

wertigkeit aller Speichen beruht. Da der Informationsgehalt jeder Speiche, der jeweils einer Projektion (»Schattenbild«) des untersuchten Objektes entspricht, in seiner Bedeutung für das zu berechnende zweidimensionale Bild gleich ist, kann auf einen großen Teil der radialen Daten verzichtet werden, ohne dass sich dies im rekonstruierten Bild durch eine dramatische Verschlechterung der Auflösung oder durch Bildfehler bemerkbar machen würde.

In einem ersten Schritt konnten wir zeigen, dass sich durch einfaches Weglassen von radialen Daten eine Beschleunigung um etwa einen Faktor 2 bis 3 erreichen lässt (Zhang et al. 2010a). Bei weiterer Datenreduktion, wie sie beispielsweise für die Untersuchung des Herzens notwendig ist, ergibt sich wie in Abbildung 3 (links) ein Bild mit sehr schlechtem Signal-zu-Rausch-Verhältnis, das zudem von starken Streifenartefakten überlagert wird. Dieses Problem konnte in unserer Arbeitsgruppe auf eine überraschend erfolgreiche Weise gelöst werden. Die Bild-Rekonstruktion durch das *Gridding*-Verfahren wurde dabei durch eine iterative Berechnung mit Methoden der numerischen Mathematik ersetzt, bei der das gewünschte Bild als die Lösung eines inversen Problems definiert wird. Die außergewöhnliche Leistungsfähigkeit des verwendeten Algorithmus zeigt Abbildung 3 (Rechts) am

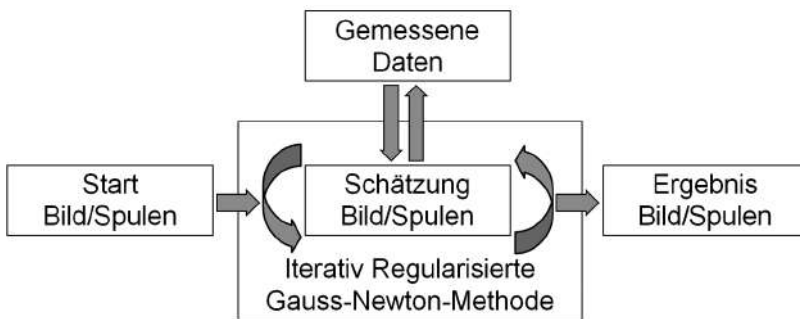


Abbildung 4: Das Prinzip der iterativen Bild-Rekonstruktionen durch regularisierte nichtlineare Inversion. Bei dieser Berechnung werden ein MRT-Bild und die Intensitätsprofile der eingesetzten Empfangsspulen geschätzt und die sich daraus ergebenden Daten berechnet. Anschließend wird das Schätzbild durch iterativen Abgleich der synthetischen mit den gemessenen Daten verändert und in mehreren Schritten optimiert. Dieser Prozess kann erheblich verbessert werden, wenn sich die Lösungsvielfalt durch eine Regularisierung des Iterationsprozesses mit Zusatzkenntnissen über das gewünschte Bild einschränken lässt.

Beispiel eines MRT-Bildes aus der besonders schnellen postsystolischen Expansionsphase des Herzens. Die beiden Bilder der Abbildung wurden einerseits mit konventionellem *Gridding* und andererseits mit der neu entwickelten Technik aus denselben 15 Speichen rekonstruiert (Messzeit: 30 Millisekunden).

Bildberechnung als inverses Problem

Die Berechnung eines MRT-Bildes erfolgt normalerweise direkt aus den Daten mittels einer festen Rechenvorschrift (Fourier-Transformation). Wie Abbildung 3 zeigt, versagt dieses Prinzip bei starker Unterabtastung des Datenraumes, die in diesem Beispiel einer Beschränkung auf nur noch 7 % der eigentlich notwendigen Datenmenge entspricht. Abbildung 4 deutet schematisch an, dass der Ansatz der inversen Lösung genau umgekehrt verfährt: Wenn das Bild nicht korrekt aus den Daten berechnet werden kann, dann schätzt man das Bild und berechnet die dazugehörigen Daten durch eine gegenüber der normalen Bildberechnung umgekehrte Rechenvorschrift. Da nach einem

beliebigen Startbild kaum die richtigen Daten an denjenigen Stellen entstehen werden, an denen gemessene Daten vorliegen, wird das Bild so lange in einem iterativen Prozess verändert, bis berechnete und gemessene Daten gut übereinstimmen. Gleichzeitig werden dabei fehlende Daten (zwischen den gemessenen radialen Speichen) ergänzt.

Für Anwendungen in der MRT ist es notwendig, außer dem Bild auch die Empfindlichkeitsprofile der vielen Empfangsantennen (Spulen) zu bestimmen, mit denen die Daten parallel aufgenommen werden. Unsere Arbeitsgruppe konnte zeigen, dass die Schätzung der richtigen Lösung am besten gelingt, wenn der Optimierungsprozess gleichzeitig für die Spulenprofile und das Bild realisiert werden kann (Uecker et al. 2008), auch wenn dies zu einem nichtlinearen inversen Problem führt, das mit einem erheblichen Rechenaufwand gelöst werden muss.

Wichtig für eine Berechnung des »richtigen« Bildes aus möglichst wenigen Daten ist dabei, dass die Optimierung mit Hilfe von Rahmenbedingungen sinnvoll eingeschränkt – regularisiert – werden kann. Für filmische Aufnahmen mit der Echtzeit-MRT, die zu einer Serie von aufeinander folgenden Bildern führen, hat sich eine zeitliche Regularisierung der aktuellen Lösung mit dem unmittelbar vorhergehenden Bild, das eine große Ähnlichkeit mit dem aktuellen Bild besitzen muss, als besonders effizient herausgestellt. Das inverse Rekonstruktionsproblem wird also so gelöst, dass das zu bestimmende Bild mit den wenigen gemessenen Daten genau übereinstimmt und aus der Vielzahl verbleibender Lösungen diejenige gefunden wird, die dem vorhergehenden Bild nahekommst.

Erste Anwendungen der Echtzeit-MRT

Zur Zeit sind die beschriebenen Ansätze der Echtzeit-MRT noch nicht auf klinisch eingesetzten, kommerziellen MRT-Geräten verfügbar. Erste Pilot-Untersuchungen befassen sich daher mit ausgewählten Erprobungen möglicher praktischer Anwendungen im Vorfeld einer späteren klinischen Nutzung. Ein Beispiel ist die anatomische und funktionelle Darstellung der Bewegung von Gelenken. Dabei stehen zunächst die normalen physiologischen Verhältnisse an gesunden Versuchspersonen im Vordergrund, bevor in einzelnen Fällen Studien an Patienten folgen können.

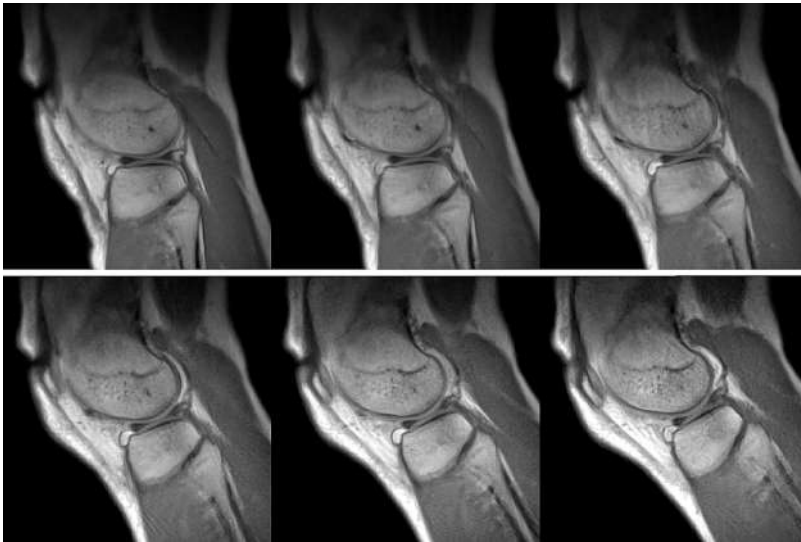


Abbildung 5: Echtzeit-MRT der Kniebewegung (ausgewählte Bilder von links oben bis rechts unten). Die Messzeit der Bilder betrug 0,33 Sekunden bei einer räumlichen Auflösung von $0,75 \times 0,75 \text{ mm}^2$ und einer Schichtdicke von 5 mm. Neben Muskelgewebe (grau), Fettgewebe (hell) und Bändern (dunkel) wird das Knie durch den Knochen bzw. das Knochenmark sowie Knorpel und Meniskus dargestellt. Außerdem zeigt die Versuchsperson nach einer Prellung kleinere Ergüsse (helle Flüssigkeitsansammlungen) im vorderen unteren und hinteren oberen Bereich des Gelenkes.

Unsere Arbeitsgruppe hat erste Erfahrungen am Fuß- und Kniegelenk gewonnen. Abbildung 5 zeigt einzelne Bilder aus einem MRT-Film, der eine langsame Beugung des rechten Knies einer Versuchsperson abbildet (in auf dem Bauch liegender Position im MRT-Magneten). Die Aufnahmen zeigen trotz einer Geschwindigkeit von nur drei Bildern pro Sekunde alle wesentlichen anatomischen Elemente des Kniegelenkes in ausreichender Bildschärfe, insbesondere neben Muskel- und Fettgewebe den Knochen (schwarz) bzw. das Knochenmark (hell), Meniskus und Knorpelstrukturen. Darüber hinaus erkennt man eine Ansammlung von Flüssigkeit in zwei Ergüssen nach einer Prellung. Ein gegenüber statischen Bildern zusätzlicher Nutzen dynamischer Aufnahmen ist allerdings erst in Verbindung mit einer Belastung des Knies zu erwarten, die dem Normalgewicht beim Laufen entspricht.

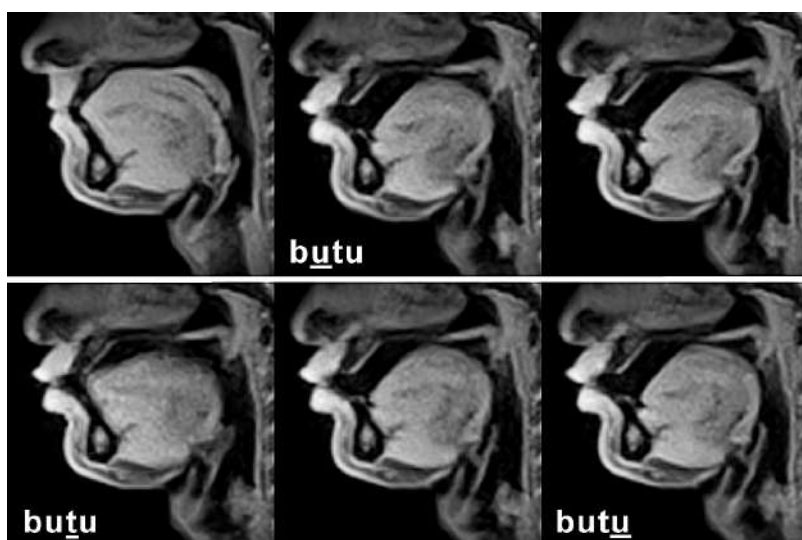


Abbildung 6: Echtzeit-MRT des Sprechens bei der Artikulation des Kunstwortes [butu] (ausgewählte Bilder von links oben bis rechts unten). Die Messzeit der Bilder betrug 42 Millisekunden bei einer räumlichen Auflösung von $1,5 \times 1,5 \text{ mm}^2$ und einer Schichtdicke von 10 mm. Die zeitliche Auflösung von 24 Bildern pro Sekunde erlaubt eine genaue Visualisierung der Stellung von Lippen, Zunge und Gaumensegel zu den jeweils durch den Unterstrich angegebenen Zeitpunkten, die den Vokalen [u] bzw. dem Konsonanten [t] zuzuordnen sind.

Dazu ist der Bau einer Vorrichtung geplant, die es erlaubt, in Rückenlage im MRT-Magneten mit dem Bein gegen eine Federkraft zu treten.

Ein unmittelbarer klinischer Nutzen der Echtzeit-MRT hat sich bereits für die Untersuchung des Kiefergelenkes ergeben. Dynamische Messungen beziehen sich dabei auf MRT-Filme des Gelenkes beim selbstständigen Öffnen und Schließen des Mundes. Bisher einzigartige und diagnostisch relevante Einblicke über die Ursachen schmerzhafter Kiefergelenksbewegungen lassen sich so in kürzester Zeit auf nicht-invasive Weise gewinnen (Zusammenarbeit mit PD Dr. N. Gersdorff, Universitätsmedizin Göttingen).

In ähnlicher Weise versprechen filmische MRT-Aufnahmen beim Sprechen und Schlucken neue Erkenntnisse. Abbildung 6 zeigt ein Beispiel aus einer systematischen Untersuchung der Bewegungen des

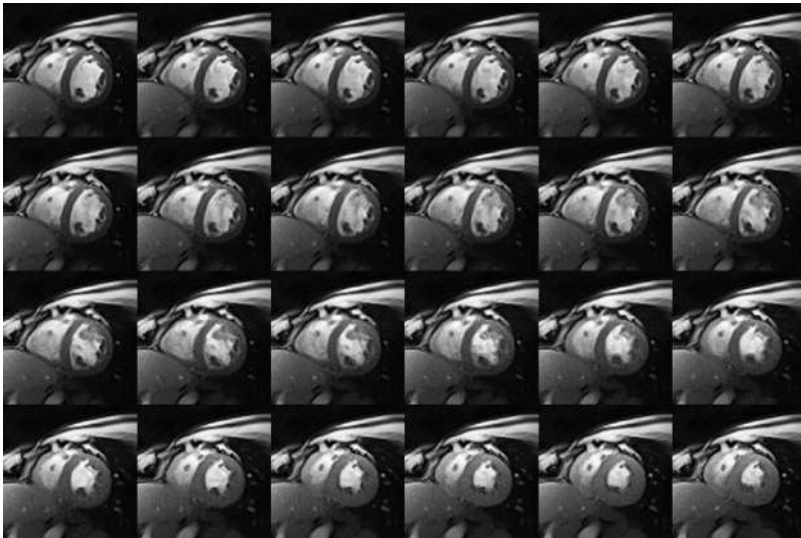


Abbildung 7: Echtzeit-MRT des Herzens einer gesunden Versuchsperson mit einer Messzeit von 30 Millisekunden pro Bild (Kurzachsenblick der rechten und linken Herzkammer). Die Abbildung zeigt 24 aufeinander folgende Bilder, die einen 0,72 Sekunden langen Ausschnitt aus einem einzelnen Herzschlag repräsentieren: Die Sequenz reicht von der Diastole (links oben) bis zur systolischen Kontraktion und Verdickung des Herzmuskels (rechts unten).

Sprechapparates bei der Artikulation von Kunstwörtern wie [butu]. Derartige Kunstwörter sind besonders geeignet, um beispielsweise die Positionen der Lippen, der Zunge und des Gaumensegels bei der Realisierung von Vokalen und Konsonanten zu untersuchen, sowohl unter normalen und pathologischen Bedingungen als auch in verschiedenen Sprachen. Bei dem ausgewählten Beispiel schnellte die Zunge beim [t] in weniger als 50 Millisekunden an den Gaumen hinter der oberen Zahnreihe. Die alternative Technik der Röntgen-Fluoroskopie kann aus ethischen Gründen nicht generell eingesetzt werden und ermöglicht auch nicht alle Schichtorientierungen, während endoskopische Beobachtungen (durch die Nase) oft die eigentlich interessierende Fragestellung stören. Dies betrifft auch die Vorgänge beim normalen Schluckakt, die in einer ersten Studie an gesunden Versuchspersonen untersucht werden (Zusammenarbeit mit PD Dr. A. Olthoff, Universitätsmedizin Göttingen).

Die neuen Möglichkeiten der Echtzeit-MRT lassen sich besonders eindrucksvoll bei der Untersuchung des menschlichen Herzens demonstrieren. Aufgrund der sehr schnellen Kontraktions- und Expansionsbewegungen des Herzmuskels ist mindestens eine zeitliche Auflösung von 50 Millisekunden erforderlich, um eine ausreichend genaue Darstellung zu erzielen. Abbildung 7 illustriert die erreichbare Qualität eines entsprechenden MRT-Filmes bei einer Messzeit von 30 Millisekunden pro Einzelbild, entsprechend einer Datenmenge von nur noch 15 Speichen (vgl. Abb. 3). Abbildung 7 enthält 24 unmittelbar aufeinander folgende Einzelbilder (0,72 Sekunden Filmausschnitt), die die Bewegungen des Herzmuskels von der Diastole (links oben) bis zur systolischen Kontraktion und Wandverdickung (rechts unten) während eines einzelnen Herzschlages darstellen.

Der Vorteil der Echtzeit-MRT des Herzens und der großen Gefäße besteht für die Patienten vor allem in der Tatsache, dass die Messung ohne das heute notwendige Anhalten des Atems erfolgen kann. Auf der klinischen Seite verbessern sich die diagnostischen Möglichkeiten um die Betrachtung einzelner Herzschläge und ihrer Variabilität. Als Stand der Technik wird bisher ein MRT-Film analysiert, der durch Synchronisation von einzelnen MRT-Messungen als Mittelwert aus 10 bis 15 Herzschlägen zusammengesetzt ist und einen einzigen synthetischen Herzzyklus repräsentiert.

Ausblick

Insgesamt ist zu erwarten, dass sich die hier nur skizzierten Möglichkeiten der Echtzeit-MRT in Zukunft erheblich erweitern. Aktuelle Fortschritte in unserer Arbeitsgruppe beziehen sich auf die Erweiterung der Herzuntersuchung um quantitative Echtzeit-Messungen der Blutflussgeschwindigkeit in den großen Herzgefäßen (Aorta, Pulmonararterie) mit einer zeitlichen Auflösung von unter 50 Millisekunden. Dazu wird die Echtzeit-Messung kombiniert mit dem Prinzip der so genannten Phasenkontrast-MRT, die zwei in Bezug auf fließende Signale unterschiedlich kodierte Einzelbilder zu einem Phasendifferenz-Bild verrechnet, dessen Signalstärke der Fließgeschwindigkeit proportional ist. Erste Anwendungen der Phasenkontrast-MRT in Echtzeit zeigen vielversprechende Ergebnisse. Dies gilt insbesondere für die Messung der Blutflussgeschwindigkeit in der aufsteigenden Aorta,

wo räumlich aufgelöste Geschwindigkeitsprofile zu verschiedenen Zeitpunkten nach Kontraktion des Herzmuskels und Öffnung der Aortenklappe einen genauen Einblick in die lokalen Strömungsverhältnisse ergeben.

Auch die so genannte interventionelle MRT wird von den hier vorgestellten Möglichkeiten profitieren. Darunter ist der Ersatz der Röntgen-Kontrolle bei minimal-invasiven Eingriffen durch eine MRT-Kontrolle zu verstehen. Voraussetzung ist allerdings eine möglichst robuste und schnelle Bildgebung in Echtzeit, die mit den neuen Verfahren in Aussicht steht.

Neben der wissenschaftlichen Weiterentwicklung und ersten Erprobungen sind zwei praktische Ziele die Anpassung der nachverarbeitenden Software (etwa von Herzfilmen) und die Reduktion der Rechenzeit. Beide Aspekte sind Voraussetzung für eine breite klinische Nutzung und Akzeptanz der Echtzeit-MRT. Dies gilt in ganz entscheidendem Maße für die algorithmische Anpassung der automatischen Segmentierung, Koregistrierung und Verrechnung der einzelnen Bilder der Echtzeit-MRT, um die qualitative Betrachtung der Herzfilme durch eine quantitative Analyse von klinischen Parametern zu ergänzen, die als Maß für die Herzleistung dienen. Dazu zählen beispielsweise die Auswurfraction des Blutes aus der linken Herzkammer und die regionale Wandverdickung des Herzmuskels (Zusammenarbeit mit Fraunhofer MEVIS, Institut für medizinische Bildverarbeitung, Bremen).

Eine ebenso wichtige Voraussetzung für die praktische Umsetzung der Echtzeit-MRT ist die Reduktion der Rechenzeit, so dass die MRT-Filme nicht nur in Echtzeit gemessen, sondern auch in Echtzeit rekonstruiert und dargestellt werden können. Im Augenblick ist eine unmittelbare Berechnung und Beobachtung während der Messung nur durch eine Kombination mehrerer aufeinander folgender Datensätze möglich, die dann mit der *Gridding*-Technik berechnet werden. Wir hoffen jedoch durch den Einsatz von Rechnern, die mit mehreren Grafikkarten für besonders schnelles paralleles Rechnen ausgestattet sind, sowie durch eine Optimierung des Algorithmus das Problem der Rechenzeit innerhalb der nächsten Monate zu lösen. Insgesamt ist mit einer breiten klinischen Erprobung der neuen Verfahren in zwei bis drei Jahren zu rechnen.

Literatur

- Frahm J., Haase A., Matthaei D., Merboldt K.D., Hänicke W. (1985) Deutsche Patentanmeldung P 3 504 734.8, 12. Februar 1985.
- Uecker M., Hohage T., Block K.T., Frahm J. (2008) Image reconstruction by regularized nonlinear inversion – Joint estimation of coil sensitivities and image content. *Magnetic Resonance in Medicine* 60, S. 674.
- Uecker M., Zhang S., Frahm J. (2010a) Nonlinear inverse reconstruction for real-time MRI of the human heart using undersampled radial FLASH. *Magnetic Resonance in Medicine* 63, S. 1456.
- Uecker M., Zhang S., Voit D., Karaus A., Merboldt K.D., Frahm J. (2010b) Real-time MRI at a resolution of 20 ms. *NMR in Biomedicine* 23, S. 986.
- Zhang S., Block K.T., Frahm J. (2010a) Magnetic resonance imaging in real time – Advances using radial FLASH. *Journal of Magnetic Resonance Imaging* 31, S. 101.
- Zhang S., Uecker M., Voit D., Merboldt K.D., Frahm J. (2010b) Real-time cardiovascular magnetic resonance at high temporal resolution: Radial FLASH with nonlinear inverse reconstruction. *Journal of Cardiovascular Magnetic Resonance* 12, S. 39.

Ulf Diederichsen

Lineare molekulare Architekturen mit DNA- und Peptid-Biooligomer-Gerüsten¹

Als die wichtigsten Vertreter der Biooligomere können die Oligonucleotide DNA und RNA einerseits und die Peptide und Proteine andererseits angesehen werden. Beiden Biooligomer-Typen ist strukturell gemeinsam, dass sie als nicht-verzweigte Ketten von aneinander gereihten Monomereinheiten anzusehen sind, die ein jeweils uniformes Rückgrat bilden. An diesem Rückgrat sorgt eine regelmäßige Sequenz von Nucleobasen bei DNA und RNA bzw. Aminosäure-Seitenketten bei den Peptiden und Proteinen für die Codierung von Information, für die spezifische Erkennung, für die dreidimensionale Strukturierung und für die Funktion beispielsweise als Enzym.

Dem Rückgrat der Oligonucleotid- und Peptidoligomere kommt wichtige Bedeutung im Hinblick auf Erkennung und Informationsweitergabe zu, da es nicht nur für eine strukturell an allen Positionen einheitliche Umgebung verantwortlich ist, die ausschließlich die sequenziell codierte Information zum Tragen kommen lässt, sondern auch einen erheblichen Beitrag leistet im Hinblick auf die Faltungsmöglichkeiten und die Rückgrattopologie. Das korrekte Auslesen der DNA-Sequenz durch komplementäre Nucleobasen-Erkennung ist beispielsweise kontrolliert durch die Topologie der Doppelhelix (Abb. 1). Sie ist entscheidend für die Eindeutigkeit der Basenkomplementarität zwischen Adenin und Thymin bzw. Guanin und Cytosin. Nur durch den in der Doppelhelix für ein Basenpaar definierten Raum und die festgelegte Orientierung der beiden Nucleobasen zueinander und zum Rückgrat ergibt sich die erforderliche Einheitlichkeit, nach der alle Basenpaarkombinationen an jedem Platz der Doppelhelix austauschbar vorkommen können. Die Abstoßung der in den Phosphodiestern lokalisierten gleichmäßig über das Rückgrat verteilten negativen Ladungen sorgt für eine gestreckte und starre DNA-Doppelhelix. Anhand der DNA-Doppelhelix lässt sich auch ein weiterer Parameter aufzeigen, der strukturell die Doppelhelix definiert: Die Stabilität der DNA-

1 Vortrag beim XLAB Science Festival 2010.

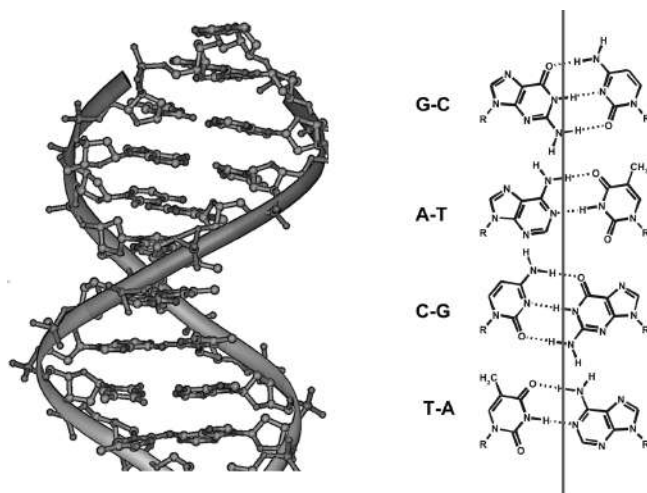


Abbildung 1: DNA-Doppelstrang (links) und Kombination der vier möglichen DNA-Basenpaarungen (rechts). In Größe und Orientierung der Ribosyl-Verknüpfung (R) aller Watson-Crick-Basenpaare ist eine Pseudosymmetrie erkennbar.

Doppelhelix basiert auf der Erkennung der Basen durch Ausbildung von Wasserstoffbrücken, aber auch auf der Interaktion der planaren aromatischen Nucleobasen durch Wechselwirkung der π -Systeme, der Basenstapelung und nicht zuletzt Solvenseffekten, bei denen das polare wässrige Lösungsmittel Wasser die unpolaren Basenpaare im Inneren der Helix aggregieren lässt. Die Basenstapelung erreicht höchste Effizienz bei einem Abstand der Basenpaarebenen von $3,4 \text{ \AA}$, so dass ein Rückgrat mit einem Abstand von drei Bindungen zwischen aufeinander folgenden Nucleobasen diesen Abstand recht genau realisieren würde. Beim Rückgrat der DNA stehen sechs Bindungen für jede Nucleotideinheit zur Verfügung, so dass das System die zusätzliche Rückgratlänge abbilden muss und dies im längeren Weg realisiert, der bei einer Helicalisierung entsteht [1]. Das System stellt den Basenpaarabstand ein, so dass die Ausbildung einer Helix fast eine Zwangsläufigkeit darstellt, um die doppelte Rückgratlänge umzusetzen. Weitere strukturelle Faktoren sind zudem verantwortlich dafür, dass das DNA-Rückgrat nicht so flexibel ist, wie man vermuten müsste, wenn sechs Bindungen mit freier Rotation vorliegen. Genannt sei zum einen der

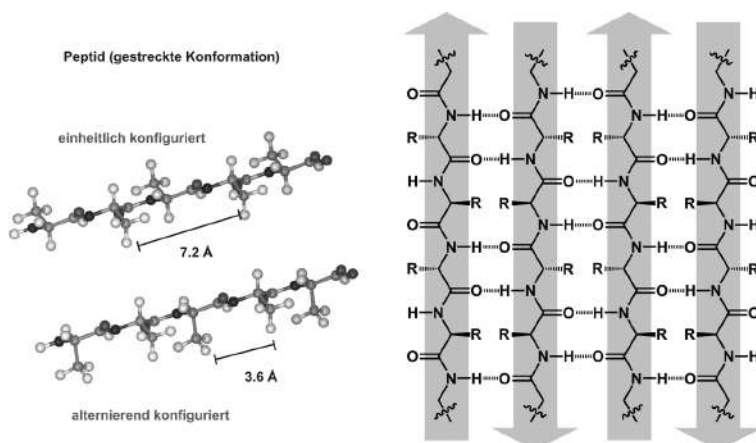


Abbildung 2: Peptidstrang in gestreckter Konformation regulär einheitlich L-konfiguriert (links oben) mit alternierender Orientierung der Methyl-Seitenketten sowie alternierend D/L-konfiguriert (links unten) mit Ausrichtung aller Methyl-Seitenketten auf einer Seite des Stranges. Organisation von gestreckten Peptidketten im antiparallelen β -Faltblatt (rechts).

im Rückgrat involvierte Fünfring, der eine erhebliche Einschränkung der Rotationsfreiheitsgrade bedeutet, sowie die Ausrichtung der den Phosphodiester ausmachenden Bindungen. Durch eine Kombination aus sterischer Hinderung (1,3-Repulsion) und einem stabilisierenden stereoelektronischen Beitrag (pseudo-anomerer Effekt, $n \rightarrow \sigma^*$ Interaktion des Elektronenpaares am Sauerstoff mit der antiperiplanaren polarisierten P-O-Bindung) ist die dreidimensionale Ausrichtung der Bindungen (Konformation) sehr gut bestimmt [2]. Die Diskussion von Basenstapelung und Festlegung der Rückgratkonformation soll verdeutlichen, dass die Rückgratstruktur der Oligomere keinesfalls zufälligen Charakter hat in Form einer Kette, an der die Erkennungseinheiten sequenziell aufgereiht sind, sondern dass für ein eindeutiges Auslesen dieser Information auch die aus der Rückgratkonformation resultierende Helixstruktur eine entscheidende Rolle spielt.

Als Beispiel für die strukturelle Bedeutung des Rückgrats bei den Peptiden soll die Sekundärstruktur des β -Faltblatts dienen (Abb. 2). Eine Peptidkette, aufgebaut aus sequenziell miteinander verknüpften Aminosäuren, stellt sich dar als eine alternierende Abfolge von Amid-

bindungen (-NH-CO-) und Methylenheiten (-CHR-). Von den drei Bindungen, die im Rückgrat des Peptids eine Aminosäure verkörpern, sind lediglich zwei frei rotierbar, da die Amidbindung aufgrund ihres Doppelbindungscharakters als planar und relativ starr anzusehen ist. Das Peptidrückgrat im β -Faltblatt ist annähernd gestreckt, so dass die Kette eine Zickzack-Ausrichtung annimmt. Derartige Ketten sind über Wasserstoffbrücken zwischen den Amideinheiten der Ketten engmaschig verknüpft, so dass über dieses Wasserstoffbrücken-Netzwerk eine Ebene aufgespannt wird. Die Aminosäure-Seitenketten (R) an den Methylenheiten sind nun im β -Faltblatt alternierend oberhalb und unterhalb der Ebene ausgerichtet, so dass dieses Strukturmotiv es beispielsweise erlaubt, eine polare Oberseite und eine unpolare Unterseite zu realisieren. Diese alternierende Ausrichtung von Seitenketten ist zum einen durch die planare Stabilisierung durch Amid-Wasserstoffbrücken gewährleistet und beruht zum anderen auf der Zickzack-Anordnung der Ketten, die zur Folge hat, dass man bei drei Bindungen eine alternierende Ausrichtung nach oben und unten erhält. Ebenfalls von Bedeutung ist die einheitliche Konfiguration (Chiralität) aller Aminosäuren. Die Ungeradzähligkeit der Bindungszahl einer Monomereinheit im Rückgrat ist demnach zusammen mit der über die Konfiguration definierten Ausrichtung der Seitenkette für das Peptid-Strukturmotiv eines β -Faltblatts von besonderer Relevanz.

Alanyl-Peptidnucleinsäuren mit linearer Doppelstrangtopologie

Mit dem Wissen um die Bedeutung von Bindungszahl, Rotationsfreiheitsgrad der Bindungen sowie Konfiguration etwaiger Stereozentren im Hinblick auf die definierte dreidimensionale Struktur von Bio-oligomeren lassen sich gezielt künstliche Oligomere mit definierter Strukturtopologie konstruieren. Ein derartiges Beispiel stellt die lineare Doppelstrangtopologie von Alanyl-Peptidnucleinsäuren (Alanyl-PNA) dar (Abb. 3), ein Hybrid aus Peptiden und Oligonucleotiden, in dem ein gestrecktes Peptidrückgrat kombiniert wird mit Wasserstoffbrücken-vermittelter Basenpaarung und Basenstapelung [3]. Ausgehend von gestreckten Peptidketten, in denen der Abstand aufeinander folgender Aminosäure-Seitenketten bei drei Rückgratbindungen etwa 3,5 Å beträgt, und der Beobachtung, dass dieser Abstand wieder

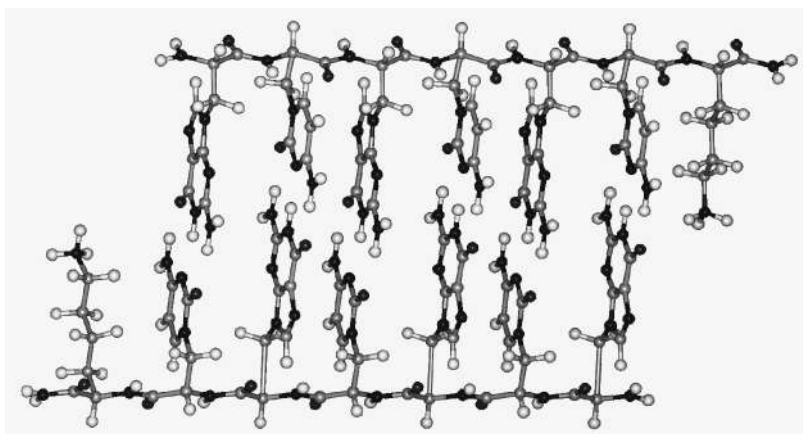


Abbildung 3: Linearer Doppelstrang einer Alanyl-Peptidnucleinsäure (Alanyl-PNA). (Bildnachweis: [3])

auftaucht als nahezu idealer Stapelungsabstand von Basenpaaren in DNA, lag es nahe, diese beiden Strukturelemente in einem Molekül zu verknüpfen, um ein gestrecktes, ideal stapelndes Doppelstrang-Paarungssystem zu generieren, das rigide, lineare Topologie aufweisen sollte. Bei der Diskussion der β -Faltblatt-Sekundärstruktur waren die in Bezug auf die Ebene der Peptidketten alternierend nach oben und unten ausgerichteten Seitenketten charakteristisch. Eine derartige Anordnung der Seitenketten (Nucleobasen) in einem PNA-Strang wäre nicht sonderlich zielführend, um mit jeder Nucleobase Erkennung des gegenüberliegenden Stranges herbeizuführen. An dieser Stelle ist die Chiralität der beteiligten Nucleoaminosäuren von Bedeutung. Da sich enantiomere Moleküle, die sich wie Bild und Spiegelbild verhalten, durch Austausch zweier an einem tetraedrischen Kohlenstoff substituierten Gruppierungen ineinander überführen lassen, bedeutet die Änderung der Konfiguration jeder übernächsten Nucleoaminosäure in Alanyl-PNA die Umorientierung dieser Seitenkette, so dass nun alle Seitenketten der gestreckten Peptidkette nahezu in die gleiche Richtung orientiert sind. Dieser Trick, alternierend D/L-konfigurierte Nucleoaminosäuren in einem Alanyl-PNA-Strang zu verwenden, erlaubt es, alle Nucleobasen, die kovalent am gestreckten peptidischen Rückgrat verknüpft sind, so zu orientieren, dass sie mit einem gegenüberliegenden Strang über Wasserstoffbrückenbindungen interagieren können.

Derartige Alanyl-PNA-Oligomere wurden synthetisch durch Herstellung der unnatürlichen Nucleoaminosäuren und anschließende Verknüpfung der Monomereinheiten zu Oligomeren mit Hilfe der Peptidfestphasensynthese erhalten. Die Nucleobasen Adenin, Guanin, Thymin und Cytosin wurden dabei über eine Methylen Einheit ($-\text{CH}_2-$) am Rückgrat kovalent unter Beachtung der D/L-Konfiguration verknüpft [3]. Von Interesse war das Vermögen von Alanyl-PNA-Oligomeren mit Strängen komplementärer Nucleobasensequenz Doppelstränge ausbilden zu können, die nach den obigen Ausführungen nahezu zwangsläufig eine lineare Doppelstrangtopologie aufweisen sollten. Tatsächlich zeigte es sich, dass stabile Alanyl-PNA-Doppelstränge bereits mit Oligomeren gebildet werden, die lediglich eine Länge von sechs Nucleoaminosäuren umfassen. Die Ausbildung der Doppelstränge war in hohem Maße abhängig von der Nucleobasensequenz, folgte allerdings nicht den von der DNA vertrauten Paarungsregeln. Neben A/T- und G/C-Paarungen wurden alle Purin/Purin- und Purin/Pyrimidin-Kombinationen realisiert. Zudem wurden neben der Watson-Crick-Paarung auch Möglichkeiten zur Erkennung über Wasserstoffbrücken über die Hoogsteen-Seite der Purine sowie Basenpaarung mit hinsichtlich der Paarung invertiert ausgerichteten Nucleobasen (inverse Watson-Crick- und inverse Hoogsteen-Paarung) verwirklicht [4]. Weder die Größe der Basenpaare noch die relative Orientierung der Nucleobasen und die in Richtung der Paarung ausgerichtete Seite sind offensichtlich in Alanyl-PNA-Doppelsträngen festgelegt. Der Vorteil der durch die Topologie definierten Eindeutigkeit der Basenkomplementarität in DNA-Doppelhelices ist bei den linearen Doppelsträngen der Alanyl-PNA offensichtlich verloren gegangen. Die größere Vielfalt an Basenpaarungserkennung in Alanyl-PNA ist aufgrund der linearen Doppelstrangtopologie verständlich. Durch die Annäherung linearer Stränge ist die Größe der die Erkennung ausmachenden Basenpaare grundsätzlich nicht festgelegt, ebenso wenig wie die relative Orientierung dieser Stränge. So ist die Verwirklichung der Vielzahl von Basenpaarkombinationen mit ganz unterschiedlichen Paarungsmodi bereits ein Indiz dafür, dass für Alanyl-PNA-Doppelstränge die lineare Doppelstrangtopologie maßgeblich sein muss. Anhand der lineare Doppelstränge ausbildenden Alanyl-PNA konnte ausgemacht werden, dass die Watson-Crick-Basenpaarung insbesondere zwischen Adenin und Thymin keinesfalls die bevorzugte Interaktionsmöglichkeit dieser Erkennungseinheiten darstellt, sondern vielmehr

mehrere gleichrangige Möglichkeiten existieren. Die Eindeutigkeit der A/T-Erkennung in DNA im Watson-Crick-Modus ist offenkundig der helicalen Doppelstrangtopologie zu verdanken. Die größere Paarungsvielfalt in linearen Alanyl-PNA-Doppelsträngen wird noch gesteigert durch den Einsatz unnatürlicher Nucleobasen, durch die sich der Buchstabencode über die vier kanonischen Nucleobasen hinaus deutlich erweitern lässt [5].

Angesichts der diskutierten unterschiedlichen Länge des Rückgrats im DNA- und Alanyl-PNA-Oligomer bei gleichem Basenpaarabstand und der daraus resultierenden Präferenz für helicale bzw. lineare Doppelstrangstruktur wird auch deutlich, dass DNA und Alanyl-PNA-Oligomere nicht miteinander in Wechselwirkung treten sollten, obwohl das Erkennungspotenzial durch Basenpaarererkennung grundsätzlich gegeben ist. Tatsächlich wurde die Eigenständigkeit dieser beiden auf Nucleobasensequenz aufbauenden Oligomere experimentell verifiziert; beide Erkennungssysteme können unabhängig voneinander und nebeneinander operieren. Es sei noch auf die ungewöhnlich hohe Stabilität von Alanyl-PNA-Doppelsträngen hingewiesen; die Stabilität eines G/C-paarenden Hexamer-Doppelstranges beträgt 58°C , verglichen mit der des entsprechenden DNA-Oligomers von 24°C . Diese Stabilisierung ist ein Resultat fehlender Abstoßung der zusammenkommenden Stränge aber auch eines hohen Grades von konformatieller Präorganisation der Einzelstränge im Hinblick auf die Duplex-Ausbildung.

Lineare Topologie auf Basis von β -Peptiden

Auf der Suche nach weiteren peptidischen Strukturen, die ein lineares Gerüst bieten, an dem Nucleobasen als Erkennungseinheiten organisiert werden können, lohnt es sich, die β -Peptide näher zu betrachten. β -Peptide sind Oligomere, die aus β -Aminosäuren aufgebaut sind, die sich von den α -Aminosäuren durch eine zusätzliche Methylengruppe ($-\text{CH}_2-$) unterscheiden, welche zur Verlängerung des Rückgrats führt [6,7]. Auf diesem Rückgrat durch Nucleobasenverknüpfung resultierende β -Homoolanyl-PNA-Oligomere (Abb. 4) bilden ebenfalls auf definierter Basenpaarererkennung beruhende Doppelstränge aus, jedoch weicht die Doppelstranggeometrie von der eines linearen Duplexes ab [8]. Dies wird verständlich angesichts eines Abstands aufeinander fol-

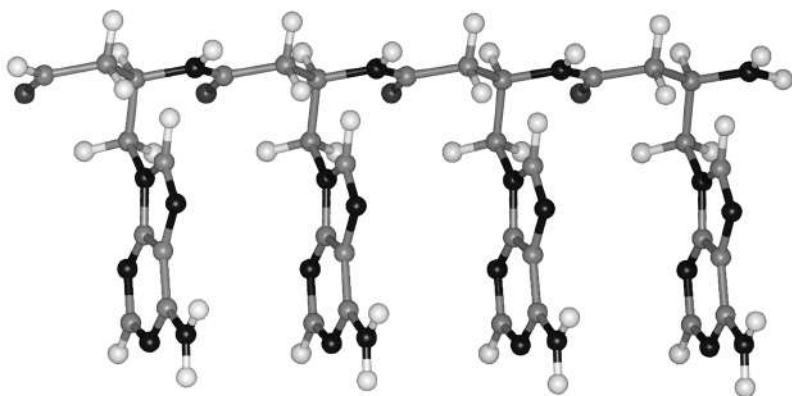


Abbildung 4: β -Homoalanyl-Peptidnucleinsäure in gestreckter Konformation. (Bildnachweis: [8])

gender Seitenketten im gestreckten β -Peptidstrang von $5 \text{ \AA}$. Dieser Abstand wird durch Helicalisierung oder Scherung verringert, so dass möglichst der ideale Basenstapelungsabstand von $3,4 \text{ \AA}$ verwirklicht wird. Dennoch bietet das Strukturmotiv einer β -Peptidnucleinsäure eine interessante Option für lineare Doppelstrangsysteme, da für β -Peptide bekannt ist, dass sie eine 1_4 -Helix ausbilden, in der bereits sechs Aminosäuren für den Erhalt einer stabilen Helix ausreichen und sequenziell jede dritte Aminosäure in die gleiche Richtung weist. Ein derart stabiles Gerüst, an dem die funktionellen Nucleobasen-Erkennungseinheiten einheitlich ausgerichtet aufgereiht werden können (Abb. 5), stellt unter dem Gesichtspunkt linearer Paarungssysteme ein topologisch interessantes System dar [9]. Anders als bei den zuvor diskutierten β -Peptidnucleinsäuren lässt sich der durch die Ganghöhe der Helix definierte idealisierte Abstand der Nucleobasenpaare von etwa $5 \text{ \AA}$ nicht durch Helicalisierung oder Scherung des Rückgrats ausgleichen, sondern es ist zu erwarten, dass die Nucleobasenenebene von der orthogonalen Orientierung gegenüber der Rückgrat-Helixachse abweicht und auf diese Weise versucht, den idealen Stapelungsabstand zu realisieren. Im Ergebnis resultiert eine durch die β -Peptid- 1_4 -Helix definierte lineare Peptidnucleinsäure-Topologie, die es erlaubt, Peptidhelices durch spezifische Basenpaarerkenntnis dreidimensional zu organisieren. Die Basenpaarerkenntnis innerhalb dieses Systems ist spezifisch und orthogonal zu den natürlichen Oligonucleotiden. Die

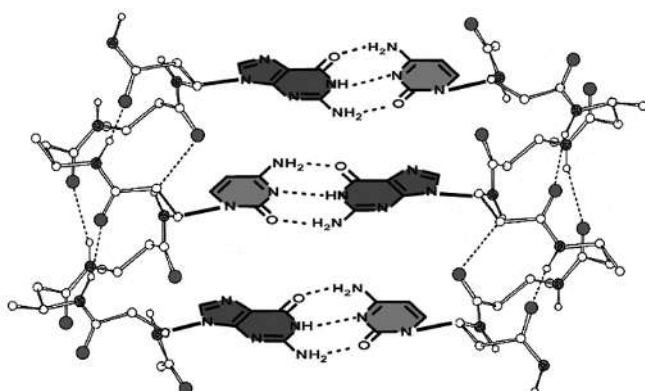


Abbildung 5: Peptidnucleinsäure-Doppelstrang mit linearem β -Peptid-14-Helix-Rückgrat. (Bildnachweis: [6,9])

β -Peptidhelices bieten einen zusätzlichen Aspekt dadurch, dass die 14-Helix drei Seiten mit gleich ausgerichteten Seitenketten aufweist. Auf diese Weise ist es bei gleichzeitiger Funktionalisierung mehrerer Flanken der Helix mit Erkennungseinheiten möglich, auch höhere Aggregate linear organisierter β -Peptidhelices zu erhalten [10].

Lineare Doppelstränge auch bei Homo-DNA

Nach der eingangs geführten Diskussion über die Zahl der Rückgratbindungen, über die aufeinander folgende Nucleobasenpaare verknüpft sind, konnte hergeleitet werden, dass durch Ausbildung einer Helix die geeignete Ganghöhe bzw. der Stapelabstand erreicht wird. Oligomere, die sich strukturell durch Aufweitung des Fünfrings um eine CH_2 -Einheit zum Sechsring von der DNA ableiten, weisen ebenso wie DNA sechs Rückgratbindungen pro Nucleotideinheit auf; es könnte demnach vermutet werden, dass sich die Homologisierung zum Sechsring nicht auf das Helicalisierungspotenzial auswirkt. Experimentell wurde der Fragestellung, warum DNA einen Desoxyribosyl-Fünfring als Zuckereinheit verwendet, nachgegangen, indem das DNA-Analogon mit einer Hexose als Zuckereinheit synthetisiert und hinsichtlich der Doppelstrang-Erkennung und Struktur untersucht wurde [11]. Dieses homologe Oligonucleotid (Homo-DNA) unterscheidet

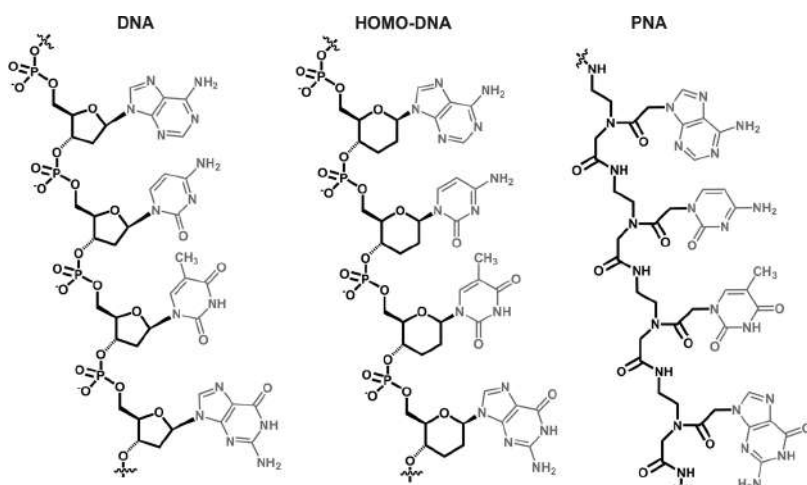


Abbildung 6: DNA-Einzelstrang (links), Homo-DNA mit homologisierten Hexose-Zuckern (Mitte) und Aminoethylglycin-Peptidnucleinsäure (rechts).

sich von natürlicher DNA ausschließlich in der Erweiterung des Desoxyribosyl-Ringes um eine Methylengruppe in allen Nucleotideinheiten (Abb. 6). Es stellte sich heraus, dass zwei Homo-DNA-Oligomere über Basenpaarung Doppelstränge ausbilden, die Homologisierung des Zuckerringes jedoch ausreicht, um von einer helicalen zur linearen Doppelstrangtopologie zu gelangen. Ganz analog zu den bereits diskutierten linearen Paarungssystemen bedeutet auch bei Homo-DNA die Linearität einen Verlust an Eindeutigkeit in der Erkennung von Nucleobasen; alleine dadurch scheidet diese DNA-Modifikation bereits als potenzieller Informationsträger für die Codierung biotischer Prozesse aus. Aus struktureller Sicht wirkt sich die Erweiterung des Fünfringes um eine Kohlenstoffeinheit in einer deutlich höheren Rigidität und Definiertheit des Zuckers aus, der in der Sesselkonformation vorliegt. In Kombination mit einem entsprechenden Substitutionsmuster am Zucker resultiert die idealisierte lineare Vorzugskonformation eines Homo-DNA-Doppelstranges (Abb. 7). Die Eigenschaften dieser Hexose-DNA lieferten einen überzeugenden Hinweis auf die Bedeutung der DNA-Rückgrat-Topologie für die Gewähr definierter Nucleobasen-Kombination und -Orientierung und damit für Eindeutigkeit im Informationstransfer.

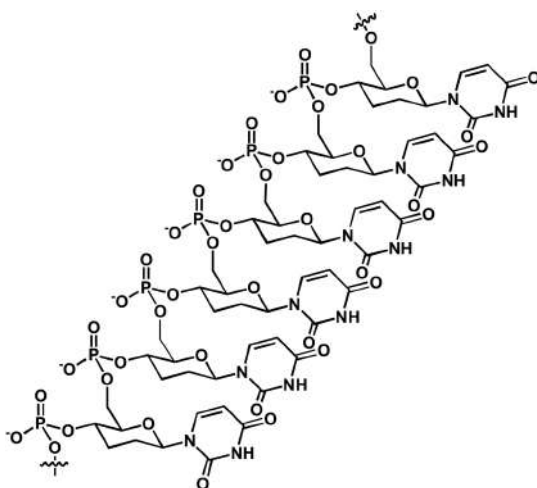


Abbildung 7: Linearer Homo-DNA Einzelstrang. (Bildnachweis: [11])

Ein auf 2-Aminoethylglycin-Einheiten basierendes Rückgrat, wie es in den Peptidnucleinsäuren (PNA) verwirklicht ist (Abb. 6, rechts), stellt insoweit das Gegenstück zur rigiden Homo-DNA dar, als in PNA ebenfalls die sechs Rückgratbindungen und eine der DNA vergleichbare Anzahl von Bindungen zur Aufhängung der Nucleobasen am Rückgrat verwirklicht sind, jedoch eine deutlich höhere Rückgratflexibilität vorliegt [12]. Als Konsequenz der höheren Flexibilität kann PNA die strukturellen Voraussetzungen einer DNA-Doppelhelix recht gut adaptieren und eignet sich daher zur Ausbildung recht stabiler und hinsichtlich der Erkennung der Basensequenz selektiver Doppelstränge.

Lineare DNA-Gerüstarchitekturen

Analog der Alanyl-PNA mit linear vorliegendem Peptidrückgrat und der β -Peptidhelix, die eine Peptidhelix verwendet, um ebenfalls eine lineare Orientierung von Erkennungseinheiten zu erlauben, verhält es sich auch mit DNA. Für das DNA-Rückgrat haben wir eindeutig helicalen Charakter dargelegt, dennoch stellt der DNA-Doppelstrang auch

ein lineares Strukturelement dar. Die Linearität ist bedingt durch die fein auf die Basenpaargeometrie und Interaktion abgestimmte Rückgratstruktur in Kombination mit der DNA als Polyanion, bei dem bereits die gleichmäßige Verteilung sich abstoßender Ladungen für eine Streckung des Stranges sorgt. Bis zu einer Persistenzlänge von 50 nm weisen DNA-Doppelstränge ein hohes Maß an linearer Rigidität auf [13]. Sowie es gelingt, den DNA-Doppelstrang an spezifischen Positionen kovalent zu funktionalisieren, liegt wiederum eine lineare Organisation dieser Funktionseinheiten vor, die gezielt in geeigneten Abständen verknüpft werden können. Als Anwendungsbeispiel der linearen Organisation von Funktionseinheiten an linearer doppelsträngiger DNA wurde kürzlich die Verwendung als Lineal zum Ausmessen von Proteinbindungstaschen berichtet [14]. Der Abstand der beiden an der DNA verknüpften Erkennungseinheiten lässt sich über die Ausbildung von DNA-Doppelsträngen durch spezifische Basenpaarerkenner nanometergenau einstellen. Als zweites Beispiel für die ortsgenaue Organisation von Funktionseinheiten in der Ebene sollen die DNA-Origami-Strukturen genannt werden. DNA-Origami lässt sich als ein planar strukturiertes System von linearen DNA-Doppelsträngen in nahezu beliebiger Form und Anordnung erhalten, indem ein langer DNA-Strang durch viele kleinere DNA-Stränge gezielt organisiert wird [15]. Über die Modifikation der kurzen DNA-Stränge lassen sich ortsspezifisch Funktionalitäten anbringen. Ein eindruckliches Beispiel für diese Strategie stellt die Generierung einer aus Proteinen bestehenden Fertigungsstraße zur Synthese kleiner Moleküle dar [16]. Verschiedene Proteine werden räumlich definiert auf der Origami-Ebene fixiert und es ist zu erwarten, dass das Substrat von einem zum nächsten Enzym übergeben wird.

Die Helicalität oder Linearität von Biooligomeren ist offensichtlich keine Zufälligkeit, sondern durch die Struktur der Rückgrateinheiten vorgegeben. Es konnten Faktoren wie z.B. die Zahl von Bindungen, Rotationsfreiheitsgrade um Einfachbindungen und Abstände von Erkennungseinheiten identifiziert werden, welche die Sekundärstruktur eines Rückgrats determinieren. Die Untersuchung künstlicher, linearer Biooligomer-Struktur motive konnte zum Verständnis beispielsweise der im Rückgrat programmierten Strukturierung der DNA-Doppelhelix und der eindeutigen Basenpaarkomplementarität beitragen. Zudem lassen sich lineare Doppelstränge ausbildende Oligonucleotide und verschiedene Peptide bzw. Peptidnucleinsäuren zum Aufbau von

Architekturen verwenden. Die Kombination von Biooligomer-Architekturen mit Funktions- oder Erkennungseinheiten eröffnet ein neues Feld und vielversprechende Optionen auf der Nanometer-Größenskala.

Literatur

- [1] G. Quinkert, E. Egert, C. Griesinger, *Aspekte der Organischen Chemie: Struktur*, Helvetica Chimica Acta, Basel 1995.
- [2] a) A. Eschenmoser, M. Dobler, *Helv. Chim. Acta* 1992, 75, 218; b) M. Böhringer, H.-J. Roth, J. Hunziker, M. Göbel, R. Krishnan, A. Giger, B. Schweizer, J. Schreiber, C. Leumann, A. Eschenmoser, *Helv. Chim. Acta* 1992, 75, 1416.
- [3] U. Diederichsen, *Angew. Chem.* 1996, 108, 458-461.
- [4] Alanyl-PNA: A model system for DNA using a linear peptide backbone as scaffold, U. Diederichsen, in *Bioorganic Chemistry – Highlights and New Aspects* hrsg. v. U. Diederichsen, T. Lindhorst, L.A. Wessjohann, B. Westermann, Wiley-VCH, Weinheim, 1999.
- [5] U. Diederichsen, D. Weicherding, N. Diezemann, *Org. Biomol. Chem.* 2005, 3, 1058-1066.
- [6] D. Seebach, A.K. Beck, D.J. Bierbaum, *Chemistry & Biodiversity* 2004, 1, 1111-1239.
- [7] R.P. Cheng, S.H. Gellman, W.F. DeGrado, *Chem. Rev.* 2001, 101, 3219-3232.
- [8] U. Diederichsen, H.W. Schmitt, *Angew. Chem.* 1998, 110, 312-315.
- [9] A.M. Brückner, P. Chakraborty, S.H. Gellman, U. Diederichsen, *Angew. Chem.* 2003, 115, 4532-4536.
- [10] P. Chakraborty, U. Diederichsen, *Chemistry Eur. J.* 2005, 11, 3207-3216.
- [11] J. Hunziker, H.-J. Roth, M. Böhringer, A. Giger, U. Diederichsen, M. Göbel, R. Krishnan, B. Jaun, C. Leumann, A. Eschenmoser, *Helv. Chim. Acta* 1993, 76, 259-352.
- [12] P.E. Nielsen, M. Egholm, R.H. Berg, O. Buchardt, *Science* 1991, 254, 1497.
- [13] J.F. Marko, S. Cocco, *Phys. World* 2003, 16, 37.
- [14] H. Eberhard, F. Diezmann, O. Seitz, *Angew. Chem.* 2011, 123, 4232-4236.
- [15] P.W.K. Rothmund, *Nature* 2006, 440, 297-302.
- [16] C.M. Niemeyer, *Angew. Chem.* 2010, 122, 1220-1238.

Hanns Ruder und Hans-Peter Nollert

Was Einstein gern gesehen hätte

Visualisierung relativistischer Effekte¹

Manchem mag die Einsteinsche Relativitätstheorie als der Inbegriff einer abstrakten, unanschaulichen Theorie erscheinen. Dabei haben sowohl die Spezielle als auch die Allgemeine Relativitätstheorie höchst anschauliche Auswirkungen.

Besser gesagt, sie hätten – wenn wir jeden Tag mit 90 % der Lichtgeschwindigkeit zu unserem Arbeitsplatz in der Nähe eines Schwarzen Loches pendeln würden, wobei vielleicht noch ein Wurmloch zu passieren wäre. Dies ist heute leider genauso wenig möglich wie zu Zeiten Einsteins.

Zum Glück gibt es inzwischen aber leistungsfähige Computer, die es uns wenigstens erlauben, die visuellen Effekte sichtbar und sogar erfahrbar zu machen. Als sichtbarer Effekt der Speziellen Relativitätstheorie fällt einem als Erstes wohl die Längenkontraktion ein: Ein bewegtes Objekt erscheint in Bewegungsrichtung um den Faktor

$$\sqrt{1 - \frac{v^2}{c^2}} \quad (1)$$

verkürzt; hier ist v die Geschwindigkeit des Objekts relativ zum Beobachter, c bezeichnet die Lichtgeschwindigkeit im Vakuum. Senkrecht zur Bewegungsrichtung bleiben alle Längen unverändert.

Man muss nicht allzu kühn sein, um zu vermuten, dass man die Längenkontraktion nicht nur messen, sondern bei sehr hohen Geschwindigkeiten auch unmittelbar sehen kann. Abbildung 1 zeigt die Erde mit einem Teleskop aus der Entfernung der Mondbahn betrachtet. Abbildung 1 a zeigt das gewohnte Bild, das entsteht, wenn Erde und Betrachter zueinander ruhen; es verändert sich nur unmerklich, solange die Geschwindigkeit einer Relativbewegung sehr viel niedriger als die Lichtgeschwindigkeit ist. Abbildung 1 b zeigt die längenkontrahierte Erde, wie wir sie beobachten könnten, wenn die Lorentz-

1 Vortrag beim XLAB Science Festival 2011.

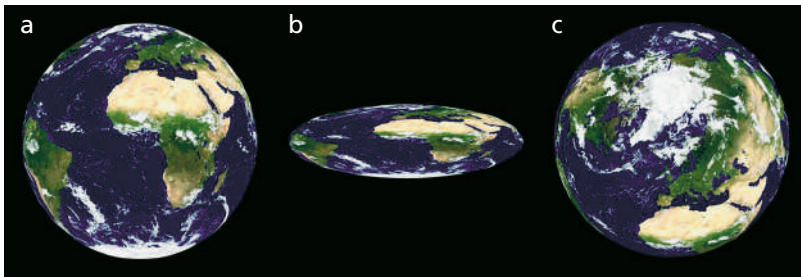


Abbildung 1: Die Erde, mit einem starken Teleskop aus der Entfernung der Mondumlaufbahn betrachtet.

- a Das gewohnte Bild, der Beobachter bewegt sich relativ zur Erde nicht.
- b Die längenkontrahierte Erde, wie wir sie sehen könnten, wenn die Lorentzkontraktion unmittelbar sichtbar wäre. Die Bewegung verläuft in Richtung der Polachse mit einer Geschwindigkeit von 95 % der Lichtgeschwindigkeit.
- c Die Erde, wie sie sich tatsächlich zeigt, wenn der Beobachter sich mit 95 % der Lichtgeschwindigkeit an ihr vorbeibewegt. Die Auswirkungen der endlichen Lichtlaufzeit kompensieren weitgehend den Einfluss der Längenkontraktion auf das Bild, so dass der Umriss der Erde wieder genau kreisförmig erscheint.

kontraktion unmittelbar sichtbar wäre. Die Bewegung verläuft in Richtung der Polachse mit einer Geschwindigkeit von 95 % der Lichtgeschwindigkeit.

Dabei kommt es nur auf die relative Bewegung an: Das Relativitätsprinzip besagt, dass es keinen absoluten Bezug für Bewegung an sich gibt; es kommt immer nur auf die relative Bewegung von Objekten oder Bezugssystemen an. Dieses Relativitätsprinzip war auch schon durch die Galilei-Invarianz der klassischen Mechanik realisiert. Als neues Element der Speziellen Relativitätstheorie kommt hinzu, dass die Geschwindigkeit eines Lichtstrahls im Vakuum – genauer gesagt: jeder physikalischen Erscheinung, die keine Ruhemasse aufweist und die sich im Vakuum ausbreiten kann – in jedem Bezugssystem die gleiche ist, auch wenn sich verschiedene Bezugssysteme ihrerseits gegeneinander bewegen. Geschwindigkeiten können daher nicht mehr, wie gewohnt, einfach zueinander addiert werden.

In Abbildung 1 b spielt es dank des Relativitätsprinzips keine Rolle, ob wir uns vorstellen, dass die Erde am Beobachter vorbeifliegt oder umgekehrt. Abbildung 1 b ist jedoch fiktiv: Ein realistisches Bild kommt selbst durch Lichtstrahlen zustande. Dabei spielt ein sehr ele-

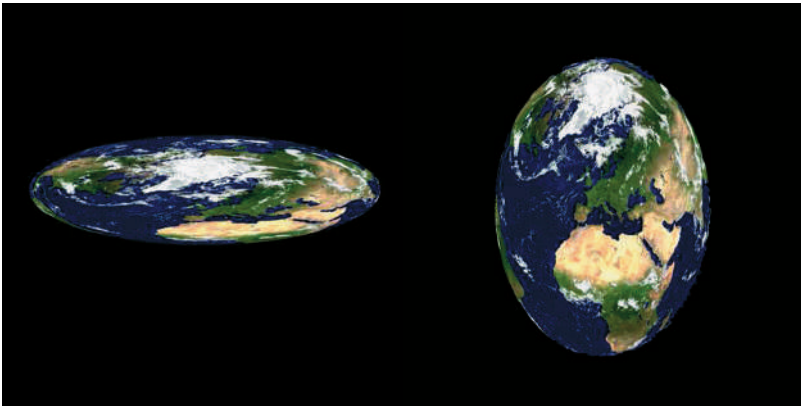


Abbildung 2. Links: Dieses Bild der Erde hätte ein Physiker vor Entwicklung der Speziellen Relativitätstheorie erwartet, und zwar für den Fall, dass die Erde im absoluten Raum ruht und der Beobachter sich mit 95 % der Lichtgeschwindigkeit relativ dazu bewegt. Rechts: Der gleiche Physiker müsste dieses Bild vorhersagen für den Fall, dass die Erde sich gegen den Äther bewegt und der Beobachter ruht.

mentarer Effekt, nämlich die endliche Geschwindigkeit des Lichts, eine wichtige Rolle. Abbildung 1c berücksichtigt die dadurch bedingte endliche Lichtlaufzeit und zeigt daher, was wir bei der genannten Geschwindigkeit tatsächlich sehen. Das Resultat erstaunt: Man hätte wohl erwartet, dass das Bild noch kurioser aussieht als Abbildung 1b. In Wirklichkeit kompensiert der Effekt der Lichtlaufzeit aber einen erheblichen Teil der Veränderungen, die die Längenkontraktion verursacht hat. Der Umriss der Erde ist wieder genau kreisförmig, auf den ersten Blick erkennt man hauptsächlich eine Drehung des gewohnten Bildes.

So simpel dieser Effekt also ist, so hat die Endlichkeit der Lichtgeschwindigkeit dennoch einen genauso großen Einfluss auf das endgültige Bild wie die Längenkontraktion. Spätestens seit den Arbeiten von Ole Rømer (z.B. Rømer 1676) ist aber bekannt, dass Licht sich mit endlicher Geschwindigkeit ausbreitet; sogar den Wert dieser Geschwindigkeit hatte Rømer recht gut bestimmen können.

Hat sich also jemals ein Physiker vor Einstein überlegt, welche Auswirkungen das bei einem Blick aus dem Fenster während einer sehr schnellen Reise hat? Anscheinend nicht. Dabei wäre er nämlich auf

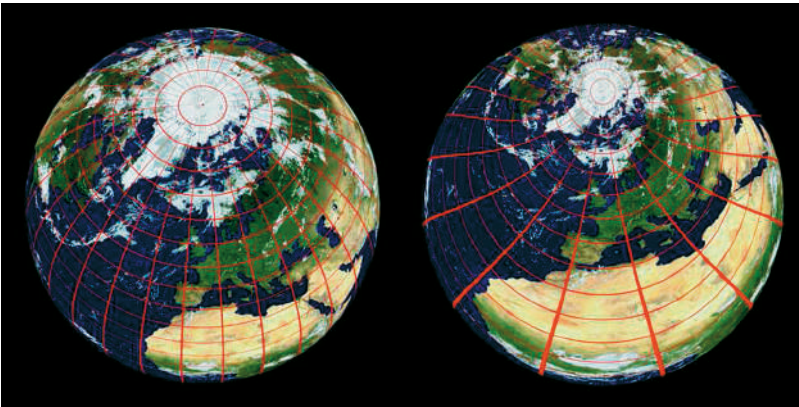


Abbildung 3. Links: Zur Verdeutlichung der Verzerrung bei relativistischer Bewegung ist hier ein Netz von Längen- und Breitenkreisen über die Erde gelegt. In diesem Referenzbild ruht der Beobachter relativ zur Erde; der Standpunkt liegt diesmal $\frac{1}{3}$ des Erdradius über der Erdoberfläche. Rechts: Hier bewegt sich der Betrachter mit 99,9 % der Lichtgeschwindigkeit relativ zur Erde. Der Abstand zur Erdoberfläche ist der gleiche wie links.

interessante Ergebnisse gestoßen: Entsprechend dem Stand der Physik vor 1900 muss er sich zunächst entscheiden, ob er die Maxwell-Gleichungen im Bezugssystem der Erde oder im Bezugssystem des Beobachters lösen will. Die Maxwellgleichungen sind die Grundgleichungen der Elektrodynamik und bestimmen daher auch die Ausbreitung des Lichts als elektromagnetische Welle. In der Sprache der Physik des 19. Jahrhunderts kann man diese Unterscheidung auch als die Frage formulieren, ob die Erde sich relativ zum Äther bewegt, oder aber der Beobachter, oder möglicherweise beide.

Das Ergebnis ist unterschiedlich je nachdem, ob wir die Erde oder den Beobachter als bewegt ansehen. So etwas galt bei den Physikern auch schon vor der Einsteinschen Relativitätstheorie als problematisch. Die Gleichungen der Newtonschen Mechanik sind nämlich invariant gegenüber einem Wechsel von einem Inertialsystem zu einem anderen. Auch hier sollte also nur die relative Bewegung eine Rolle beim Ausgang eines Experiments spielen.

Was wir in Abbildung 2 anschaulich dargestellt sehen, ist eines der Hauptprobleme der Physik im ausgehenden 19. Jahrhundert: Die

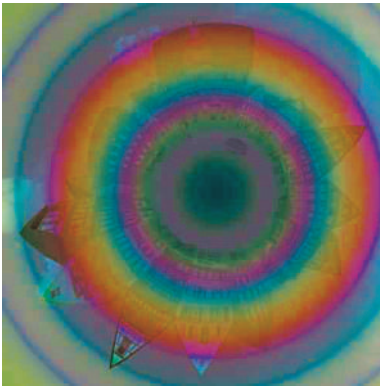


Abbildung 4: Dopplereffekt.

Maxwellgleichungen gehorchen einer anderen Art von Invarianz als die Newtonsche Mechanik. Die beiden sind also gewissermaßen nicht miteinander verträglich. Nicht von ungefähr gab Einstein seiner grundlegenden Arbeit zur Speziellen Relativitätstheorie den Titel »Zur Elektrodynamik bewegter Körper«.

Kehren wir also ins beginnende 21. Jahrhundert zurück, in dem diese Probleme gelöst sind, die Spezielle Relativitätstheorie hervorragend experimentell bestätigt ist und ihren festen Platz

in der Physik hat. Wir haben in Abbildung 1 c bereits gesehen, dass der Umriss der Erde kreisförmig bleibt, auch wenn sich der Betrachter mit hoher Geschwindigkeit an ihr vorbei bewegt; die Längenkontraktion wird durch den Lichtlaufzeiteffekt in dieser Hinsicht ausgeglichen. Dennoch wird das Bild gegenüber dem Fall ohne relative Bewegung verändert, was in Abbildung 1 c nicht so deutlich zu erkennen ist. In Abbildung 3 wurde daher ein Netz aus Längen- und Breitenkreisen über die Erde gelegt; die Verzerrung ist hier deutlich zu sehen.

Bisher haben wir nur die geometrischen Veränderungen gezeigt, die das Bild eines schnell bewegten Gegenstandes erfährt. Zwei weitere Effekte sind der Dopplereffekt und der *Searchlight*-Effekt. Der Dopplereffekt ist uns in seiner akustischen Variante vertraut: Nähert sich ein Einsatzfahrzeug mit großer Geschwindigkeit, ist der Ton seiner Sirene zu höheren Frequenzen verschoben, entfernt es sich wieder, wird der Ton sehr viel tiefer. Das gleiche passiert mit der Frequenz des Lichts, das von einem Objekt stammt, das sich mit relativistischer Geschwindigkeit auf uns zu bewegt: Vor allem in der Mitte des Bildes erkennen wir eine starke Verschiebung zu blauen Farben hin. Entfernt sich das Objekt von uns, sehen wir es dagegen stark nach Rot verschoben (Abb. 4).

Der *Searchlight*-Effekt beruht teilweise auf der Aberration der Lichtstrahlen. Aberration bedeutet, dass Lichtstrahlen bei schneller Bewegung bevorzugt von vorne zu kommen scheinen. Auch dieser Effekt hat eine Analogie im Alltag: Fährt man schnell durch fallenden

Schnee, so scheinen die Schneeflocken hauptsächlich von vorne zu kommen, obwohl sie tatsächlich (bei Windstille) gleichmäßig von oben nach unten sinken. Da die Lichtstrahlen also gewissermaßen gebündelt werden, ist der unmittelbar vor einem liegende Teil des Bildes erheblich heller als seitliche Teile, oder gar das Bild bei einem Blick nach hinten (Abb. 5).



Abbildung 5: Searchlight-Effekt.

Nicht nur der Raum, auch die Zeit ist bei Einstein nicht mehr, was sie bei Newton einmal war. Was die Längenkontraktion für den Raum, das ist die Zeitdilatation für die Zeitdimension der vierdimensionalen Raumzeit: Der Zeitablauf in einem bewegten System erscheint um den gleichen Faktor verlangsamt (Gleichung (1)). Dies sieht man sehr schön an Myonen, die in ca. 10 km Höhe durch die kosmische Höhenstrahlung erzeugt werden (Abb. 6). Diese Myonen bewegen sich fast mit Lichtgeschwindigkeit. Da sie eine mittlere Lebensdauer von ca. $2,2 \mu\text{s}$ ($1 \mu\text{s}$ ist ein Millionstel einer Sekunde) haben, würde man vielleicht erwarten, dass sie im Mittel ca. 660 m weit kommen, bevor sie wieder zerfallen. Trotzdem weist man aber noch am Erdboden eine relativ große Zahl von Myonen nach. Tatsächlich ist die Flugzeit des Myons, würde sie im Bezugssystem der Erde gemessen, erheblich länger; im bewegten Bezugssystem des Myons dagegen ist dank der Zeitdilatation nur so viel Zeit vergangen, wie es der Lebensdauer des Myons entspricht.

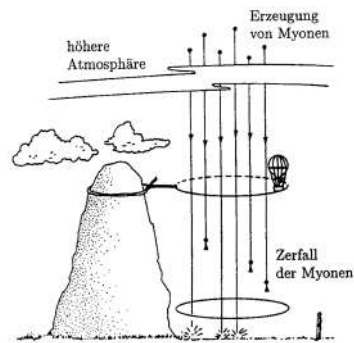


Abbildung 6: Das Myonenexperiment von Rossi und Hall, 1941. (Details des Experiments in Hoffmann 1997).



Abbildung 7: Ausblick beim senkrechten Sturz auf den Tübinger Marktplatz mit 90 % der Lichtgeschwindigkeit.

Dennoch wird im Bezugssystem des Myons die gleiche Geschwindigkeit relativ zur Erde gemessen. Zwar ist die verstrichene Zeitspanne kürzer; vom System des Myons aus ist jedoch die Erde bewegt, so dass die Entfernung vom Entstehungsort zum Erdboden längenkontrahiert ist und damit um den gleichen Faktor kürzer erscheint wie die verstrichene Zeit.

Dabei ist wichtig, dass man die Bezugssysteme nicht unzulässig kombiniert; dies führt unweigerlich zu widersprüchlichen Ergebnissen. Bestimmt man etwa im obigen Beispiel die Länge der

vom Myon zurückgelegten Strecke im Ruhesystem der Erde (wie man das normalerweise tun wird) und nimmt die Lebensdauer als Obergrenze für die Flugzeit, so folgert man irrtümlich, dass das Myon erheblich schneller als mit Lichtgeschwindigkeit unterwegs war. Ein gleichermaßen falsches Ergebnis erhält man, wenn man für einen Raumflug beispielsweise zum Zentrum unserer Milchstraße die von der Erde aus gemessene Entfernung, ca. 30 000 Lichtjahre, durch die im Raumschiff gemessene Reisezeit dividiert.

Abbildung 7 zeigt den Tübinger Marktplatz von oben aus der Sicht eines Myons, kurz bevor es mit knapp Lichtgeschwindigkeit den Erdboden erreicht.

Allgemeine Relativitätstheorie

Bisher haben wir uns mit den visuellen Effekten der Speziellen Relativitätstheorie beschäftigt. Auch die Allgemeine Relativitätstheorie hält eine erstaunliche Fülle von optisch eindrucksvollen Kuriositäten bereit. Eine Bewegung des Objekts relativ zum Beobachter ist dabei nicht erforderlich. Ursache der Verzerrungen ist vielmehr die relativistische Lichtablenkung, also die Tatsache, dass sich Photonen nicht mehr auf

Geraden im Raum bewegen. Im Rahmen der Allgemeinen Relativitätstheorie werden nämlich nicht nur die Bahnen massiver Körper durch andere Massen beeinflusst, sondern auch die Bahnen von Photonen, d.h. also Lichtstrahlen. Formal kommt dies daher, dass sich Photonen, ebenso wie massive Körper, auf so genannten Geodäten der Raumzeit bewegen. In einer gekrümmten Raumzeit spielen Geodäten die gleiche Rolle wie Geraden in einer euklidischen Raumzeit (also einem flachen Raum mit einer unabhängigen, überall gleich ablaufenden Zeit): Sie sind die kürzesten Verbindungen zwischen zwei Punkten. Die Krümmung der Raumzeit wird ihrerseits von allen vorhandenen Massen erzeugt. Anschaulich kann man diese Lichtablenkung verstehen, indem man sich an die berühmte Formel $E=mc^2$ erinnert: Da Photonen Energie besitzen, kann man ihnen eine Masse (wenn auch keine Ruhemasse) zuordnen. Auf diesem Wege kann man sogar im Rahmen der Newtonschen Gravitationstheorie die Bahn eines Photons berechnen. Das Ergebnis ist qualitativ ähnlich, als Ablenkung kommt aber nur die Hälfte des korrekten relativistischen Wertes heraus. Der Unterschied lässt sich auf die allgemein relativistische Zeitdilatation zurückführen: In der Nähe großer Massen vergeht die Zeit langsamer als weit davon entfernt.

Historisch lieferte die Ablenkung von Lichtstrahlen, die nahe an der Sonne vorbeilaufen, eine der ersten Möglichkeiten, die Vorhersagen der Allgemeinen Relativitätstheorie zu überprüfen. Die Gelegenheit dazu bot sich bei der Sonnenfinsternis von 1919; der theoretisch vorhergesagte Wert von 1,75 Bogensekunden konnte eindrucksvoll bestätigt werden. Mittlerweile lässt sich mit hochgenauen Radioteleskopmessungen sogar die Lichtablenkung am Rand der Erde nachweisen, obwohl diese nur 0,575 Tausendstel einer Bogensekunde beträgt!

Mit bloßem Auge lassen sich solche geringen Ablenkungen natürlich nicht erkennen. Deutlich sichtbar wird der Effekt aber, wenn wir die Ablenkung des Lichts an einem sehr kompakten Objekt, etwa einem Schwarzen Loch oder einem Neutronenstern, betrachten.

Schauen wir also wieder auf unsere wohlbekannte Erde, wobei wir ein Schwarzes Loch – mit ausreichendem Sicherheitsabstand – vor der Erde vorbeiziehen lassen. Nähert sich das Schwarze Loch, wie in der Sequenz in Abbildung 8 gezeigt, der Sichtlinie, wird das Bild zunächst in ähnlicher Weise verzerrt, wie dies der Fall wäre, wenn sich eine Linse zwischen Erde und Betrachter befände. Nicht von ungefähr spricht man ja von Gravitationslinsen – wobei nicht nur Schwarze Löcher, sondern auch Sterne, Galaxien oder ganze Galaxienhaufen als Gravita-

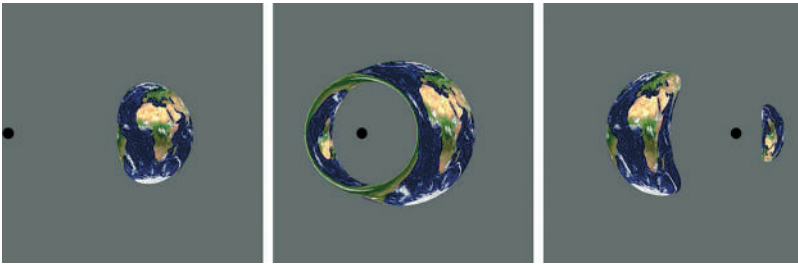


Abbildung 8: Ein Schwarzes Loch durchquert den Bereich zwischen der Erde und dem Betrachter. Der Betrachter ist ca. 830 000 km von der Erde entfernt, also knapp das Dreifache der Entfernung zwischen Erde und Mond. Das Schwarze Loch zieht mit einem Abstand von ca. 240 000 km an der Erde vorbei, seine Masse beträgt das 326fache der Masse unserer Sonne. Ist das Schwarze Loch nahe genug an der Sichtlinie, dann wird die Ablenkung stark genug, so dass Lichtstrahlen nicht nur rechts, sondern auch links um das Schwarze Loch herum den Betrachter erreichen können. Es erscheint daher ein zweites Abbild der Erde auf der anderen Seite des Schwarzen Loches. Dieses zweite Abbild ist spiegelbildlich zum Hauptbild; man erkennt dies besonders schön am Umriss von Afrika ganz rechts.

tionslinsen wirken können. Befindet sich das Schwarze Loch schließlich genau auf der Linie zwischen Erde und Betrachter, so entsteht eine ringförmige Struktur, da nun Lichtstrahlen auf allen Seiten um das Schwarze Loch herum zum Betrachter gelenkt werden können. Abbildung 9 zeigt das Phänomen des Einsteinrings für Schwarze Löcher mit unterschiedlichen Massen bei gleichen Abständen zwischen Erde, Schwarzem Loch und Betrachter.

Nicht nur die Erde rotiert um ihre eigene Achse, alle Sterne im Universum, selbst Schwarze Löcher, tun dies ebenso. In der Newtonschen Gravitationstheorie hat die Rotation keinen Einfluss auf das Gravitationsfeld, das von einem Stern erzeugt wird. Anders in der Einsteinschen Relativitätstheorie: Hier entsteht ein Mitführungseffekt, der so genannte *Lense-Thirring*-Effekt, der die Bahnen von massiven Objekten wie auch von Photonen beeinflusst. Je nach Rotationsgeschwindigkeit und -richtung entsteht also beim Betrachter ein anderes Bild als ohne Rotation, wenn die Lichtstrahlen, die das Bild aufbauen, nahe an einem kompakten rotierenden Objekt vorbeigelaufen sind. Abbildung 10a zeigt erneut einen Einsteinring, diesmal verursacht durch ein rotierendes Schwarzes Loch. Der Unterschied zum nicht rotierenden

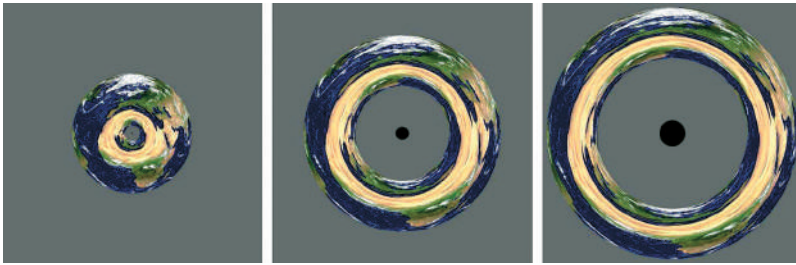


Abbildung 9: Ein Schwarzes Loch befindet sich genau auf der Linie zwischen der Erde und dem Betrachter. Infolge der Symmetrie dieser Situation können Lichtstrahlen auf allen Seiten um das Schwarze Loch herum gelenkt werden. Die entstehende ringförmige Struktur bezeichnet man als Einsteinring. Im linken Bild beträgt die Masse des Schwarzen Lochs das 16,3fache der Masse unserer Sonne, im mittleren Bild das 163fache und im rechten das 326fache. Die Abstände sind die gleichen wie in Abbildung 8.

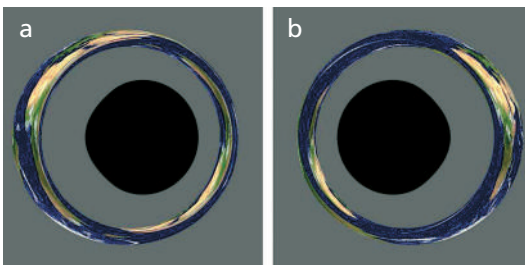


Abbildung 10: Schwarzes Loch in Rotation.

a Hier wirkt ein rotierendes Schwarzes Loch als Gravitationslinse und erzeugt einen Einsteinring, der infolge der Rotation asymmetrisch erscheint. Das Schwarze Loch rotiert mit der maximalen Geschwindigkeit, die die Relativitätstheorie erlaubt.

b Hier rotiert das Schwarze Loch entgegengesetzt zu Abbildung 10 a.

Fall ist recht gering; am ehesten erkennt man ihn im Vergleich mit Abbildung 10 b; hier rotiert das Schwarze Loch mit der gleichen Geschwindigkeit, aber entgegengesetzt zu dem von Abbildung 10 a.

Da ein Schwarzes Loch selbst kein Licht emittiert, kann man es nicht »sehen«: Beim Beobachter kommen nur Lichtstrahlen an, die von anderen Objekten emittiert worden und am Schwarzen Loch vorbeigelaufen sind, aber keine, die den Horizont durchquert hätten, also aus

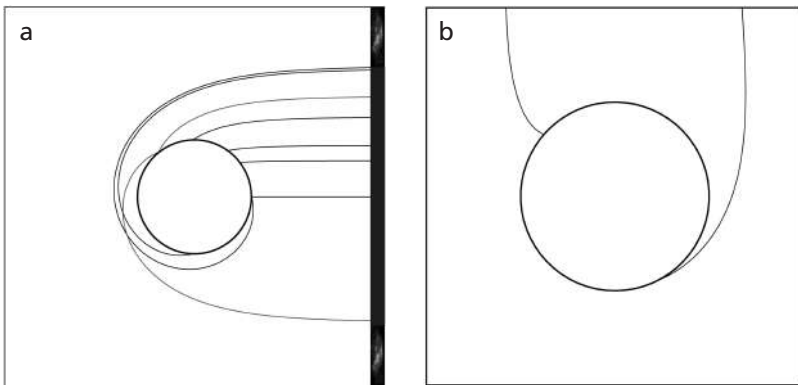


Abbildung 11: Lichtstrahlen in der Umgebung eines sehr kompakten Neutronensterns in Ruhe (a) und in Rotation im mathematisch negativen Sinn, d.h. im Uhrzeigersinn (b).

dem Loch selbst herausgekommen wären. Bei einem Neutronenstern ist dies anders: Lichtstrahlen können sehr wohl von der Oberfläche des Sterns emittiert werden und einen entfernten Beobachter erreichen. Dies ist zwar astrophysikalisch sehr interessant, visuell aber eher langweilig, da die Oberfläche eines Neutronensterns kaum strukturiert ist. Wir verabschieden uns daher von der real existierenden Erde als Anschauungsobjekt und wenden uns einer fiktiven Erde zu, die zwar noch die bekannten Kontinente zeigt, dabei aber Masse und Radius eines Neutronensterns aufweist. Ein realer Neutronenstern entsteht als eine mögliche Existenzform eines Sterns am Ende seiner Entwicklung; er stellt die kompakteste Form eines Objekts aus Materie dar, die in der Natur möglich ist (Novak 2003). Ein solcher Neutronenstern weist ungefähr das 1,4fache der Masse unserer Sonne auf; diese Masse ist aber auf eine Kugel von nur ca. 20 km Durchmesser komprimiert! An der Oberfläche eines solchen extremen Objekts herrscht eine enorm starke Gravitation, die dafür sorgt, dass die Oberfläche praktisch völlig glatt ist, da jegliche Unebenheit sofort von der übermächtigen Gravitation »geplättet« wird. Ein Neutronenstern weist also keine Struktur auf, die sich für eine Visualisierung eignen würde. Wir versehen daher einen solchen Neutronenstern künstlich mit der vertrauten Oberflächenstruktur der Erde – gewissermaßen ein Neutronenstern, der in der Haut der Erde steckt. Wie man an der Zeichnung in Abbildung 11 a erkennen kann, werden die Lichtstrahlen durch die Lichtablenkung

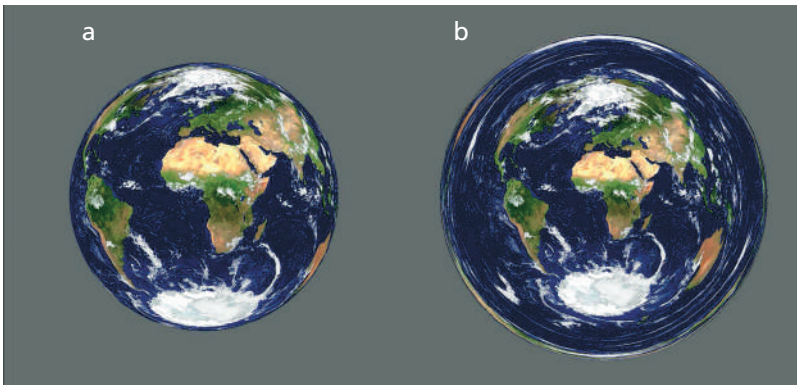


Abbildung 12: Kompakte Masse.

a Hier sehen wir einen Neutronenstern mit einer Masse von ca. 1,4 Sonnenmassen, wie er in der Natur tatsächlich vorkommt. Dieser Neutronenstern wurde allerdings als Erde »verkleidet«, um die Effekte der Lichtablenkung besser sichtbar zu machen. Da die Lichtstrahlen teilweise um den Stern herum gebogen werden, sind auch Teile der Sternoberfläche zu sehen, die dem Betrachter anderenfalls verborgen blieben. Besonders deutlich sieht man dies am Beispiel von Australien, das nun am rechten unteren Rand sichtbar wird, ohne Lichtablenkung (Abb. 1 a) dagegen verdeckt war.

b Hier sehen wir ein fiktives Objekt mit der gleichen Masse wie oben, aber einem Durchmesser von nur gut 9 km. Damit handelt es sich um das kompakteste Objekt, das gemäß der Allgemeinen Relativitätstheorie zulässig ist, ohne dass der Druck im Zentrum unendlich groß wird. In der Natur kann ein solches Objekt jedoch nicht existieren, da es keine Materieform gibt, die seine stabile Existenz ermöglichen würde. Die Lichtablenkung wird hier noch stärker als im linken Teilbild: Es ist nun die gesamte Oberfläche sichtbar, und zwar sogar mehrfach. So ist Australien nun einmal am rechten unteren und einmal am linken oberen Rand zu erkennen.

gewissermaßen um den Stern herum gelenkt. Es erreichen daher auch Strahlen von Teilen der Rückseite des Sterns den Beobachter – ohne Lichtablenkung wäre dies nicht möglich. So ist in Abbildung 12 a etwa Australien zu sehen, das auf dem Referenzbild (Abb. 1 a) nicht sichtbar war. In Abbildung 12 b treiben wir dieses Spiel noch ein wenig weiter. Hier ist die Masse des Sterns so weit erhöht und der Radius gleichzeitig so weit verringert worden, wie es im Rahmen der Allgemeinen Relativitätstheorie überhaupt möglich ist, ohne dass im Inneren des Sterns ein unendlich hoher Druck entsteht. In der Natur kann ein solcher Stern jedoch nicht existieren, da es keine realistische Materieform gibt,

die eine stabile Existenz eines solchen Objekts erlaubt. Abbildung 11 a zeigt, dass Lichtstrahlen den Stern nun komplett umrunden können; ebenso können Lichtstrahlen von der Rückseite des Sterns sowohl »rechts« als auch »links« herum den Betrachter erreichen. Dies wird in Abbildung 12 b wiederum an Australien deutlich, das einerseits rechts unten und andererseits, stark verzerrt, links oben am Rand zu erkennen ist.

Science Fiction, relativistisch

Neutronensterne und Schwarze Löcher mögen uns als fantastische Gebilde erscheinen, und doch ist inzwischen so gut wie sicher, dass sie in unserem Universum existieren – und nicht einmal als besondere Ausnahmeerscheinung. Darüber hinaus haben sich Science-Fiction-Autoren immer wieder von der Einsteinschen Relativitätstheorie zu noch ausgefalleneren Phantasien inspirieren lassen, manchmal mehr, manchmal weniger in Übereinstimmung mit der korrekten Theorie. Zwei dieser Konstrukte wollen wir nun betrachten: Wurmlöcher und Warp-Antrieb. Beide dürften zwar auch in ferner Zukunft nicht realisierbar sein, sie sind aber durch die Allgemeine Relativitätstheorie zumindest nicht grundsätzlich ausgeschlossen.

Will man als Science-Fiction-Autor eine Zivilisation beschreiben, die sich über einen größeren Bereich erstreckt als beispielsweise ein Sonnensystem, stößt man sofort auf das Problem der großen Entfernungen und damit der langen Zeitspannen, die für Reisen wie auch die Übertragung von Nachrichten benötigt werden. Mit Zeitspannen im Bereich von Wochen sind zwar die alten Römer durchaus klargekommen, hier sprechen wir jedoch von Jahren bis Jahrtausenden!

Dies gilt zunächst nur für die Daheimgebliebenen: Für die Reisenden im Raumschiff ist eine Reise zum Zentrum unserer Milchstraße und zurück dank der speziell-relativistischen Zeitdilatation grundsätzlich durchaus in gut 40 Jahren zu bewältigen – währenddessen sind aber auf der Erde ungefähr 60 000 Jahre vergangen! Dieses Synchronizitätsproblem dürfte den Aufbau einer interstellaren Zivilisation vor große organisatorische Probleme stellen – man stelle sich nur einmal die Auswirkungen auf ein Rentenversicherungssystem vor!

Kein Grund zum Verzagen, die Allgemeine Relativitätstheorie hält Alternativen bereit! Sehr beliebt im Katalog futuristischer Reiseveran-

stalter sind Wurmlöcher, die als eine Art Abkürzung zwischen weit entfernten Punkten unseres Universums dienen oder sogar zwei ansonsten vollkommen getrennte Universen verbinden können (Davies 2003). Der Begriff Wurmloch ist dabei leider etwas irreführend: Ein Wurmloch ist nämlich keine röhrenartige Konstruktion. Es besteht eher aus zwei kugelförmigen Gebilden – eines in jeder der beiden Welten – die von außen zunächst die gleiche räumliche Geometrie aufweisen wie ein Schwarzes Loch. Diese Gebilde hängen so ineinander, dass man, wenn man sich auf das vermeintliche Schwarze Loch zubewegt, irgendwann nicht mehr weiter nach innen gelangt, sondern sich vielmehr in der anderen Welt wieder nach außen bewegt. Die beiden Umgebungen um das Wurmloch herum können dabei an verschiedenen Orten im gleichen Universum liegen, oder sie können sich in zwei ansonsten völlig getrennten Universen befinden. Das Durchqueren eines Wurmlochs erfordert keine besonders lange Zeit, und auch außerhalb des Wurmlochs vergeht nicht mehr Zeit als für den Reisenden (die Details sind allerdings von der genauen Struktur des Wurmlochs abhängig; so könnten Wurmlöcher möglicherweise sogar für eine Reise in die Vergangenheit dienen). Man wird hier also nicht nur das Problem der langen Reisezeit, sondern auch das der Synchronizität los!

Die Abbildungen 13 und 14 zeigen ein solches Wurmloch mitten auf dem Tübinger Marktplatz; das andere Ende befindet sich im Weltall, von der Erde genauso weit entfernt wie der Mond. Das Bild vermittelt einen guten Eindruck von der kugelförmigen Gestalt des Wurmloches. Man kann dieses Wurmloch aus allen Raumrichtungen gleichermaßen »betreten«. Dies erkennt man noch einmal deutlich an der Sequenz Abbildung 13 unten, bei der wir ein Stück um das Wurmloch herum gehen; die Position des Wurmlochs selbst verändert sich dabei nicht. Das letzte Bild zeigt das Wurmloch von oben; außer Pflastersteinen sieht man hier zwar nicht viel, man erkennt aber wieder deutlich, dass das Wurmloch nicht eine in irgendeine bestimmte Richtung führende Röhre ist.

Das Wurmloch weist in seiner Umgebung die gleiche räumliche Struktur wie ein Schwarzes Loch auf, daher werden Lichtstrahlen in ähnlicher Weise »verbogen« wie in der Umgebung eines Schwarzen Loches. Der geneigte Leser ist daher sicher auch nicht mehr sonderlich überrascht, dass er in einem ringförmigen Bereich um das Wurmloch herum wieder das Bild seines gesamten Universums sieht. Anders als beim Schwarzen Loch gibt es hier aber keinen Horizont, vielmehr

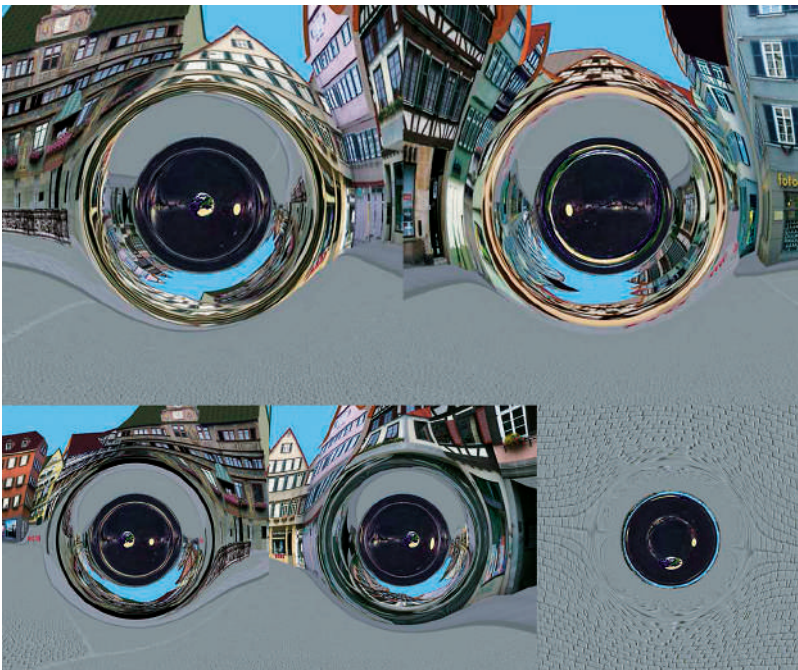


Abbildung 13: Ein Wurmloch auf dem Tübinger Marktplatz. Oben rechts sehen wir links am Wurmloch vorbei einen Teil des Rathauses. Oben links stehen wir zwischen Rathaus und Wurmloch und blicken in die entgegengesetzte Richtung. Dadurch ändert sich in gleicher Weise die Blickrichtung im inneren Teil, wo der Blick von der anderen Seite des Wurmlochs ins Universum hinausgeht: Oben rechts blicken wir auf die Erde, dahinter ist die Milchstraße zu sehen. Oben links schauen wir von der Erde weg, sehen die Erde aber noch als Ring, der das gesamte Universum einschließt. Untere Reihe: Wir gehen um das Wurmloch herum und betrachten es aus zwei weiteren Richtungen, im linken Bild direkt von oben. Das Wurmloch verändert dabei seine Position nicht. Es gibt keine »Rückseite«; aus allen Richtungen erkennt man die kugelförmige Gestalt.

sieht man im inneren Bereich die Umgebung auf der anderen Seite des Wurmlochs. Auch hier hat man einen umfassenden, nichts auslassenden Rundumblick!

Wurmlöcher eignen sich zwar gut für Reisen zu vorher bestimmten Punkten – ungeeignet sind sie aber für den Individualreisenden, der seine Reiseroute spontan entscheidet, oder gar für eine Forschungs-



Abbildung 14: Wir haben das Wurmloch durchquert und befinden uns nun auf der anderen Seite im Weltall. Wir haben uns herumgedreht und blicken durch das Wurmloch zurück zum Tübinger Marktplatz.

mission to boldly go where no man has gone before. Zum Glück gibt es da noch das Warp-Triebwerk: Hier wird nicht das Raumschiff relativ zu seiner Umgebung beschleunigt, vielmehr wird ein Teil dieser Umgebung mit dem Raumschiff darin innerhalb der gesamten Raumzeit verschoben. In Vorwärtsrichtung wird der Raum gestaucht, hinter dem Raumschiff gedehnt, wie man dies beispielsweise mit einem Gummistück ausprobieren könnte (Abb. 15). Hier können beliebig hohe Geschwindigkeiten erreicht werden, da die Beschränkung auf die Lichtgeschwindigkeit nur für eine Bewegung relativ zur Umgebung gilt, nicht aber für Verzerrungen der Raumzeit selbst. Die Reise benötigt zwar noch einige Zeit, die von der Geschwindigkeit der Verzerrung abhängig ist, aber immerhin vergeht die Zeit innerhalb der Warp-Blase genauso schnell wie außerhalb. Als Bonus ist die Besatzung des Raumschiffs keinen Trägheitskräften als Folge hoher Beschleunigung ausgesetzt (Details zum Warp-Triebwerk in Ford & Roman 2000).

Die technischen Probleme bei der Realisierung solcher Reisepläne sollen nicht verschwiegen werden: Bei der mittlerweile wohl eher konventionell anmutenden Variante einer speziell-relativistischen, fast lichtschnellen Reise stößt man vor allem auf das Treibstoffproblem: Bei einem Materie-Antimaterie-Triebwerk, das die denkbar günstigste

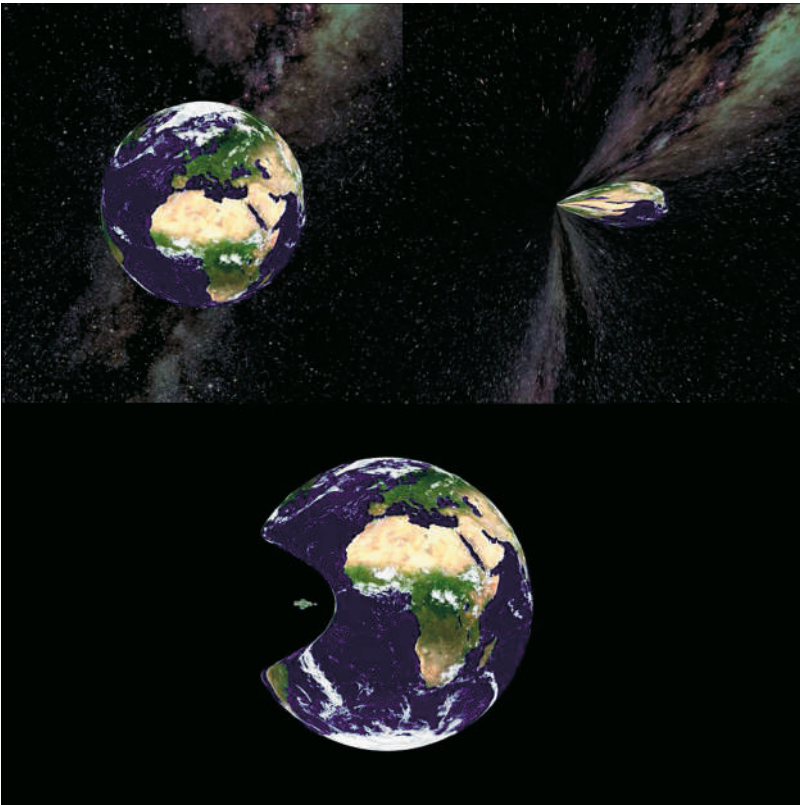


Abbildung 15. Oben: Wir befinden uns in einem Raumschiff, das von einer Warp-Blase eingeschlossen ist. Die Geschwindigkeit ist Warp 1,1. Links: Die Warp-Blase bewegt sich genau auf die Erde zu. Die Blickrichtung ist in Flugrichtung nach vorne. Rechts: Diesmal bewegt sich die Blase mit dem Raumschiff von der Erde weg, und wir schauen in Flugrichtung nach hinten auf die Erde.

Unten: Wir befinden uns außerhalb der Warp-Blase und schauen auf die Erde, während sich die Warp-Blase mitsamt einem Raumschiff darin mit Warp 1,1 zwischen uns und der Erde hindurchbewegt.

Ausnutzung des Treibstoffs erlaubt, benötigt man für die oben erwähnte Reise zum Zentrum der Milchstraße das 30000fache der Masse von Raumschiff plus Ladung allein während der Beschleunigungsphase beim Hinflug. Beim Abbremsen am Zielort ist erneut ein

entsprechender Treibstoffaufwand nötig. Man beachte, dass der beim Abbremsen benötigte Treibstoff ja zunächst mitbeschleunigt werden musste – damit braucht man insgesamt eine Menge an Materie-Antimaterie-Treibstoff, die fast eine Milliarde mal so schwer ist wie die Masse des Raumschiffs und seiner Ladung!

So gewaltig diese Anforderungen sind, so erscheinen sie doch fast banal im Vergleich zu der Herausforderung, ein stabiles Wurmloch zu betreiben, das einem Raumschiff als Abkürzung auf dem Weg zu fernen Welten dienen könnte. Es ist durchaus denkbar, dass Wurmlöcher auf natürlichem Wege entstehen, beispielsweise während des Urknalls. Ein solches natürliches Wurmloch ist jedoch nicht stabil; es zieht sich so schnell zusammen, dass ein Raumschiff keine Chance hätte, es zu durchqueren. Wer es dennoch probiert, wird unweigerlich von dem kollabierenden Wurmloch in der Singularität zerquetscht.

Andererseits bietet die Allgemeine Relativitätstheorie die Möglichkeit, ein stabiles Wurmloch zunächst zu postulieren und dann die Feldgleichungen gewissermaßen rückwärts zu benutzen, um herauszufinden, welche Verteilung von Materie man bereitstellen müsste, um die Existenz eines solchen Wurmlochs zu sichern. Das Ergebnis einer solchen Rechnung enthält aber eine unangenehme Überraschung: Erforderlich ist Materie mit negativer Energie (oder negativer Masse). Solche Materie ist in unserer Alltagswelt vollkommen unbekannt. In dem nach wie vor recht dunklen Bereich, in dem sich Relativitätstheorie und Quantentheorie treffen, kann solche Materie aber wahrscheinlich vorkommen. Dummerweise allerdings nur in Verbindung mit riesigen Mengen an Materie mit positiver Energie: Für unser gewünschtes Wurmloch wäre insgesamt ungefähr so viel Materie nötig, wie das gesamte bekannte Universum enthält! Von den technischen Schwierigkeiten, solche ungeheuren Mengen in die erforderliche Form zu bringen, wollen wir hier erst gar nicht reden ...

Der Warp-Antrieb bietet hier leider keinen Ausweg: Auch dafür ist Materie mit negativer Energie nötig. Darüber hinaus könnte ein solcher Antrieb möglicherweise gar nicht von dem Raumschiff selbst gesteuert werden, und vielleicht könnte das Raumschiff nicht einmal selbständig in die Warp-Blase eintreten oder sie verlassen. Die letzten Bemerkungen haben uns nun auf den harten Boden der Tatsachen zurückgeführt. So schade es auch sein mag, dass die von der Relativitätstheorie grundsätzlich erlaubten, phantastischen Möglichkeiten für Reisen durchs Weltall wohl nie Realität werden können – so erfreulich

ist es, dass uns moderne Computer und leistungsstarke Teleskope die Möglichkeit geben, wenigstens die visuellen Eindrücke einer solchen Reise zu genießen. Dabei ist nicht nur der erforderliche Aufwand unvergleichlich geringer, diese Möglichkeit steht zudem allen zur Verfügung und ist nicht auf die wenigen Besatzungsmitglieder eines Raumschiffs beschränkt.

Literatur

- Davies P. (2003) Bauanleitung für eine Zeitmaschine. Spektrum der Wissenschaft 03 Plus.
- Ford L.H., Roman T.A. (2000) Wurm Löcher und Überlicht-Antriebe. Spektrum der Wissenschaft 03; S. 36.
- Hoffmann B. (1997) Einsteins Ideen: Das Relativitätsprinzip und seine historischen Wurzeln. Berlin (Spektrum Akademischer Verlag).
- Novak J. (2004) Neutronensterne: Ultradichte Exoten. Spektrum der Wissenschaft 03, S. 34.
- Rømer O. (1676) *Démonstration touchant le mouvement de la lumière trouvé par M. Rømer de l'Académie royale des sciences*. Le Journal des Sçavans, Paris, S. 233.
- Ruder H., Nollert H.P. (2005) Einsteins Holodeck. Spektrum der Wissenschaft 07, S. 56.

Farbabbildungen



Farbabbildung 1: Nobelpreisträger Manfred Eigen, Schirmherr des XLAB Science Festivals, führte während seiner Laufbahn zahlreiche junge Wissenschaftler in die internationale wissenschaftliche Gemeinschaft ein (Bild: Theodoro da Silva). (Wissenschaft ist international, S. 10)



Farbabbildung 2: Wilder Tabak (*Nicotiana attenuata*, Wildtyp) an seinem natürlichen Standort im *Great Basin Desert* in Utah, USA (Foto: Celia Diezel). (Ökologische Gentechnik, S. 60)



Farbabbildung 3: Nachtschatten (*Solanum nigrum*, Wildtyp) auf einem Versuchsfeld nahe Dornburg in Thüringen, Deutschland (Foto: Markus Hartl). (Ökologische Gentechnik, S. 68)



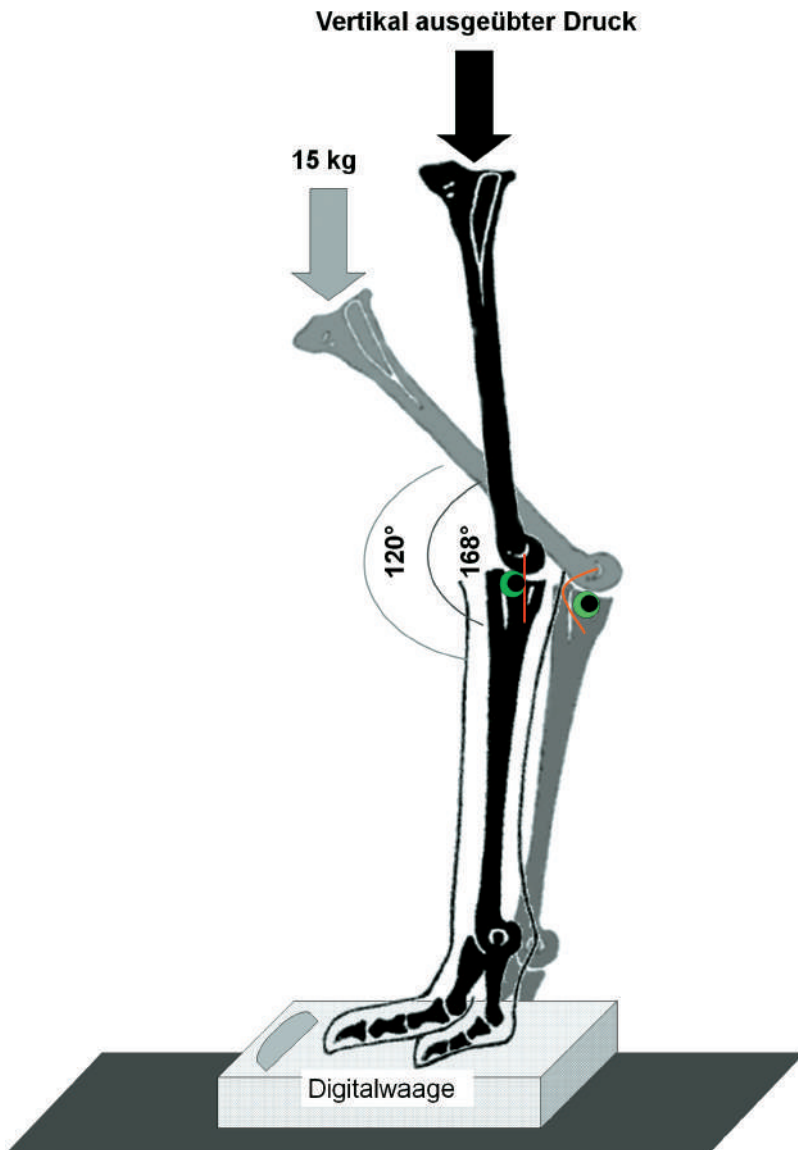
Farbabbildung 4: Larve des Tabakschwärmers (*Manduca sexta*) auf einem Blatt des Wilden Tabaks (Foto: Danny Kessler). (Ökologische Gentechnik, S. 61)



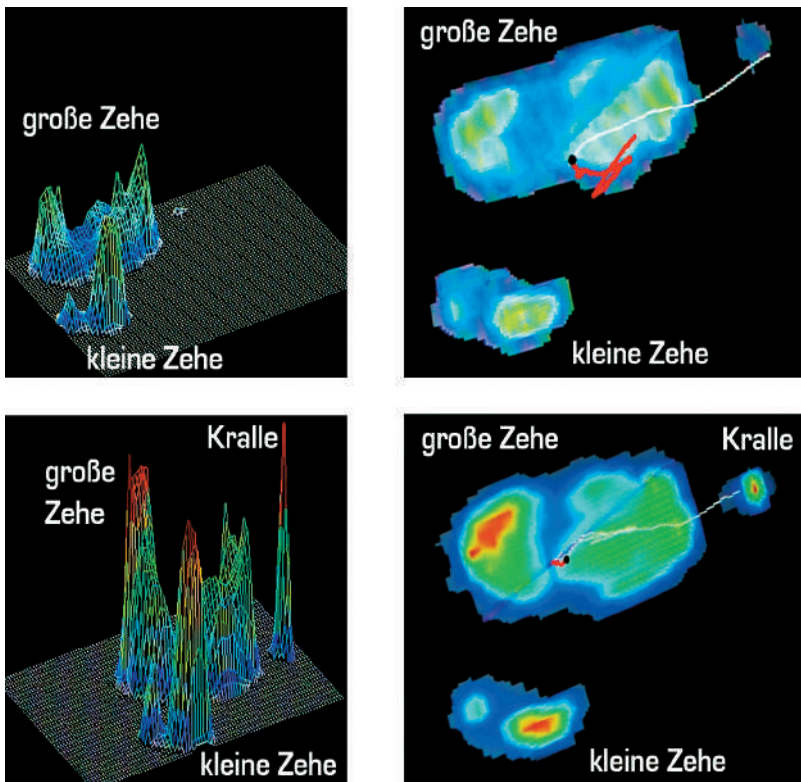
Farbabbildung 5: Eine Zwergzikaden-Art der Gattung *Empoasca*, die im Feld bevorzugt auf transgenem Wildem Tabak auftrat, bei dem ein Lipoxxygenase-Gen mittels RNAi stillgelegt wurde (Foto: André Kessler). (Ökologische Gentechnik, S. 67)



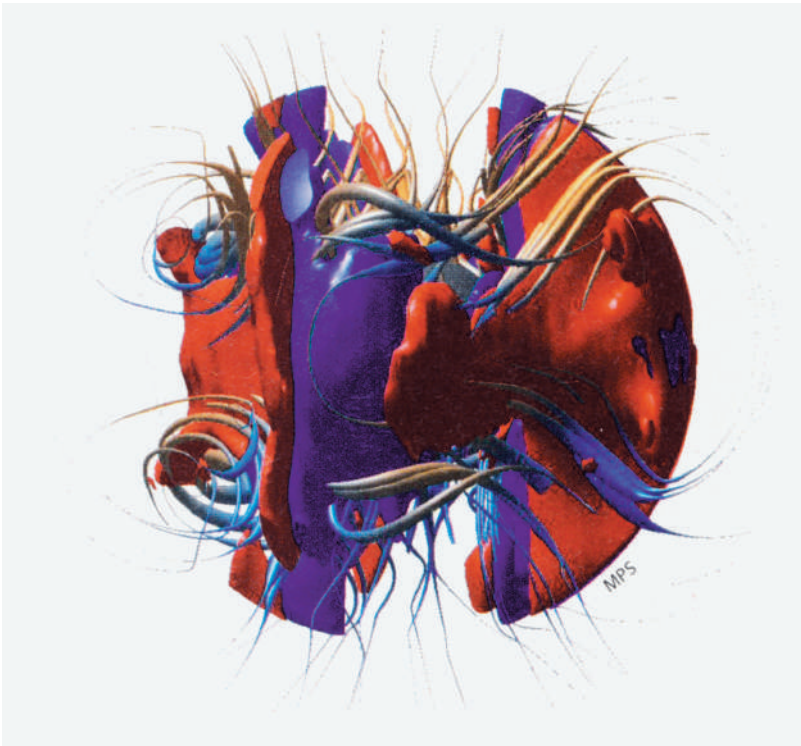
Farbabbildung 6: Flohkäfer (*Epitrix pubescens*) auf einem Blatt des Schwarzen Nachtschattens mit den für den Käfer charakteristischen kreisrunden Fraßstellen (Foto: Danny Kessler). (Ökologische Gentechnik, S. 72)



Farbabbildung 7: Das durch den Einrastmechanismus stabilisierte seziierte Straußenbein (grün: Knochenzapfen; rote Linie: Seitenband des Fersengelenks) kann ein Gewicht von bis zu 15 kg tragen. (Auf Zehenspitzen zum Weltrekord, S. 170)



Farbabbildung 8: Druckverteilung der beiden Straußenzehen (rechter Fuß) beim Gehen (oben) und Rennen (unten) in drei- und zweidimensionaler Ansicht. Beim Gehen wird die Kralle der großen Zehe kaum belastet, während beim Rennen sehr hohe Drucke auf ihr lasten: Die Kralle funktioniert bei schnellem Rennen wie ein Positionsanker (unten). Typisch für alle Zweibeiner ist, dass beim Rennen höhere Drucke auftreten, als beim Gehen. (Bildnachweis: Schaller). (Auf Zehenspitzen zum Weltrekord, S. 172)



Farbabbildung 9: Diese Dynamosimulation zeigt die komplexe Struktur magnetischer Feldlinien im flüssigen Eisenkern eines Planeten (Bildnachweis: MPS). (Die Geschwister der Erde, S. 158)

Ulrich Christensen und Norbert Krupp

Die Geschwister der Erde¹

Mit Beginn des Raumfahrtzeitalters vor 50 Jahren hat unsere Kenntnis über die Planeten des Sonnensystems größere Sprünge gemacht als durch die Einführung des Teleskops vor 400 Jahren. Inzwischen haben zahlreiche Sonden die Planeten umflogen, diese zum Teil aus einer Umlaufbahn eingehend studiert oder sie sind wie bei Mars, Venus, dem Erdmond und dem Saturnmond Titan sogar weich gelandet. Ist alles Wichtige nun bekannt und geht es bei weiteren Raummissionen nur noch darum, weiße Flecken auf der Landkarte zu füllen? Keinesfalls, denn die Kette von teilweise völlig unerwarteten Entdeckungen ist in den letzten zehn Jahren nicht abgerissen.

Die Planetenforschung liegt in den Händen von Wissenschaftlern, die sich primär mit der Erde befassen, wie z.B. Geologen und Mineralogen, Geophysiker, Plasmaphysiker, Meteorologen und sogar Geobiologen. Das wichtigste Ziel dabei ist es zu ergründen, welche Gemeinsamkeiten und Unterschiede zwischen der Erde und ihren planetaren Geschwistern bestehen und insbesondere wie und weshalb sich ihre Entwicklungswege unterscheiden. Wir wollen verstehen, was die Erde besonders macht, warum sich auf ihr lebensfreundliche Bedingungen eingestellt haben und ob es solche Bedingungen auch auf anderen Planeten gegeben hat oder sogar noch gibt.

Für die Denker der Antike bestand die irdische Materie aus den vier Elementen Feuer, Wasser, Luft und Erde, während die Himmelskörper aus etwas Reinem, dem Äther, aufgebaut waren. Mit dem Teleskop entdeckten Astronomen, dass auch andere feste Himmelskörper meist aus profanem Gestein (»Erde«) bestehen und erkannten schließlich, dass auch die übrigen »Elemente« bei allen Planeten eine wichtige Rolle spielen. Vulkanismus (»Feuer«) ist Ausdruck dafür, dass das Innere eines Planeten nicht kalt und starr ist, sondern dass unter der Oberfläche konvektive Strömungsprozesse ablaufen. Die Vulkan-

¹ Ulrich Christensen und Norbert Krupp: Die Geschwister der Erde. Physik Journal 8 (2009), Nr. 5, S. 31-36. Copyright Wiley-VCH Verlag GmbH & Co. KGaA. Reproduced with permission.

eruptionen erneuern die Oberfläche des Planeten und sind eine Quelle für atmosphärische Gase und Wasser an der Oberfläche. Eine Atmosphäre (»Luft«) dämpft die sonst extremen Temperaturschwankungen an der Oberfläche und ist Voraussetzung dafür, dass dort Flüssigkeiten existieren können. Wasser, vor allem in flüssiger Form, ist nach unserem heutigen Verständnis unabdingbar für die Entwicklung und die Existenz von Leben. In diesem Artikel wollen wir ein besonderes Augenmerk auf das Vorkommen dieser »Elemente« und der mit ihnen verbundenen dynamischen Rolle auf anderen Planeten unseres Sonnensystems richten.

Mit der neuen Definition des Begriffs »Planet« wurde Pluto zu einem Zwergplaneten degradiert. Für Planetologen ist die Unterscheidung ziemlich irrelevant. Sie bezeichnen jeden genügend großen Körper, der annähernd Kugelgestalt hat, als Planeten – selbst wenn er als Satellit einen anderen Planeten umkreist. Die Kugelgestalt impliziert stets, dass das Innere des Himmelskörpers in seine verschiedenen Komponenten (z.B. Metall, Gestein oder Eis) differenziert ist und dass der Körper zumindest in seiner Frühzeit im Inneren aktiv war und seine Oberfläche umgestaltet hat.

Wenn Planeten Feuer spucken

Aktiver Vulkanismus ist der augenfälligste Beleg dafür, dass das Innere eines Planeten dynamisch aktiv ist. Direkt beobachtet wurde vulkanische Aktivität außer bei der Erde bisher beim Jupitermond Io, dessen Inneres durch die Gezeitenreibung des nahen Jupiters aufgeheizt wird, und – völlig unerwartet – beim kleinen, eisigen Saturnmond Enceladus. Auf dem Mars sehen wir riesige Schildvulkane (Abb. 1).

Auf der Erde entsteht der Vulkanismus auf drei verschiedene Arten. Zwei davon hängen mit der Plattentektonik zusammen, die es auf dem Mars nicht (mehr) gibt und die daher als Ursache ausscheiden. Dort kommen nur so genannte Mantelplumes infrage – säulenförmige Ströme von heißem Silikatgestein, die von der Grenze zwischen Gesteinsmantel und Eisenkern aufsteigen. Wenn das heiße Gestein flache Tiefen erreicht, in denen Druck und Schmelztemperatur geringer sind, schmilzt es zu 5 bis 30 % auf. Das entstehende basaltische Magma bahnt sich einen Weg an die Oberfläche. Auf der Erde gibt es mindestens 30 solcher auf Mantelplumes beruhende Vulkanregionen, z.B. auf



Abbildung 1: Der Schildvulkan Olympus Mons, aufgenommen von der Raumsonde *Mars Express*, ist mit 24 Kilometern Höhe die größte Erhebung in unserem Sonnensystem (Bildnachweis: ESA/DLR/FU Berlin, G. Neukum).

Hawaii. Auf dem Mars finden sich die meisten großen Vulkane in einer begrenzten Region, der Tharsis-Aufwölbung. Möglicherweise gibt es aufgrund der Struktur des Marsmantels dort nur ein oder zwei Plumes. Wegen der geringeren Schwerkraft auf dem Mars treten druckinduzierte Phasenumwandlungen der Silikatmineralien, die bei der Erde im oberen Drittel des Mantels liegen, erst knapp oberhalb der Grenze zum Eisenkern auf, also in der mutmaßlichen Quellenregion von Plumes. Diese Phasenumwandlungen behindern die Konvektionsbewegungen des Mantels, und Computersimulationen legen nahe, dass sie im Mars dazu führen können, dass sich die heißen Aufströmungen in nur ein bis zwei großen Plumes konzentrieren (Harder und Chris-

tensen 1996). Die Tharsis-Vulkane sind bereits Milliarden von Jahren alt – ist der Vulkanismus auf dem Mars heute erloschen?

Kraterstatistiken dienen dazu, das Alter einer Planetenoberfläche zu bestimmen, von der es keine direkten Proben gibt. Mithilfe der Mondoberfläche, deren Alter dank der Apollo-Proben bekannt ist, lässt sich diese Methode kalibrieren. Krater entstehen durch kleine und mittelgroße Körper, die regellos und über längere Zeit in einem gleichmäßigen Strom die Oberfläche bombardieren. Eine Region, die mit vielen Kratern übersät ist, ist demnach älter als eine mit wenigen. Kleine und sehr junge Oberflächenregionen sind jedoch schwierig zu datieren, da sie nur wenige Krater aufweisen. Wichtig ist es daher, auch sehr kleine Krater zu erkennen, die naturgemäß viel häufiger vorkommen als große Vertiefungen. Dabei spielt die Auflösung der verwendeten Kameras eine entscheidende Rolle.

Mit den Bildern einer hoch auflösenden Kamera an Bord der europäischen Raumsonde *Mars Express* (mit 2,5–10 m pro Pixel) ließ sich das Alter der großen Caldera-Region des Olympus Mons zu 100 bis 200 Millionen Jahre ermitteln – relativ wenig verglichen mit dem Planetenalter von 4,5 Milliarden Jahren. Einige Flankenregionen von Olympus Mons sind sogar nur 2,5 Millionen Jahre alt (Neukum et al. 2004). Der Vulkanismus auf dem Mars ist somit geologisch gesehen erst vor kurzem erloschen, selbst wenn die vulkanische Aktivität in der jüngeren geologischen Vergangenheit episodisch und von langen Ruhephasen unterbrochen sein könnte.

Auch auf unserem Schwesterplaneten Venus, deren Oberfläche unter einer geschlossenen Wolkendecke verborgen liegt, hat die amerikanische Magellan-Sonde durch Radarkartierung zahlreiche Vulkanstrukturen auf der Oberfläche entdeckt. Da die Venus fast so groß ist wie die Erde und daher im Inneren vermutlich noch aktiver als der deutlich kleinere Mars ist, bestehen hier größere Hoffnungen, einen Vulkanausbruch *in flagranti* beobachten zu können. Im nahen Infrarot ist die Venusatmosphäre durchscheinend, sodass sich ein ausgedehnter heißer Lavastrom oder Magmasee als zeitlich variabler heller Punkt in den Infrarotinstrumenten auf der europäischen Mission Venus Express bemerkbar machen, die seit 2006 unseren Nachbarplaneten umkreist. Bisher wurde nichts gefunden, aber kleinere Ausbrüche würden wegen der starken Streuung des Signals in der Atmosphäre verborgen bleiben.

Bei der Erde trägt der Vulkanismus große Mengen von Gasen wie Schwefeldioxid in die Atmosphäre ein. In der Venusatmosphäre sind

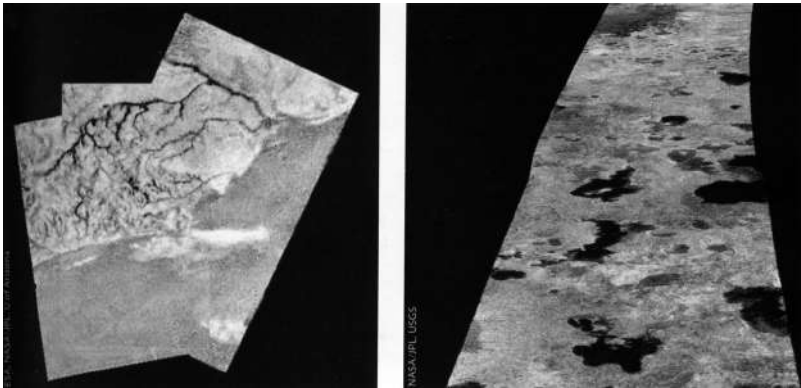


Abbildung 2: Links: Auf dem Saturnmond Titan finden sich Flussläufe, die mit flüssigem Methan gefüllt sind (Bildnachweis: ESA, NASA/JPL, University of Arizona). Rechts: Die Cassini-Radaraufnahmen zeigen ganze Methan-Seen auf der Nordhalbkugel (Bildnachweis: NASA/JPL, USGS).

recht hohe und vor allem im Laufe von Monaten und Jahren stark schwankende Konzentrationen von SO_2 zu beobachten. Schwefeldioxid sollte aber nicht lange in der Atmosphäre verbleiben, da es durch photochemische Umwandlung in Schwefelsäure (aus der die Wolken der Venusatmosphäre bestehen) und durch chemische Reaktion mit den Oberflächengesteinen wieder abgebaut wird. Die hohen Konzentrationen und starken Schwankungen könnten somit darauf hindeuten, dass aktive vulkanische Prozesse SO_2 episodisch in die Atmosphäre eintragen.

Auf der Erde entstehen durch den Vulkanismus magmatische Krustengesteine. Zudem sind große Teile des Ozeanwassers und der Atmosphäre dabei an die Oberfläche gekommen (zum Teil können sie auch von Kometen herrühren, die mit der Erde kollidiert sind). Auf den eisigen Monden des äußeren Sonnensystems spielt sich Ähnliches mit veränderten Rollen ab. Der Saturnmond Titan ist von einer dichten Atmosphäre und einer kaum durchsichtigen Smogsschicht umgeben. Auf Radarbildern der Raumsonde Cassini sind Strukturen zu erkennen, die Schildvulkanen ähneln (Abb. 2), bei denen es sich aber um sogenannten Kryovulkanismus handelt: Kryovulkane speien keine glutflüssige Lava, sondern leicht schmelzbare Substanzen wie Methan, Kohlendioxid, Wasser oder Ammoniak, die im Inneren des Planeten



Abbildung 3: Links: Bei ihrem Vorbeiflug hat die Raumsonde Cassini auf dem Saturnmond Enceladus aktiven Kryovulkanismus entdeckt (hier als *artist view* gezeigt). Rechts: Aus Schluchten steigen regelrecht Geyshire auf (Bildnachweise: NASA/JPL).

oder Mondes in gefrorenem Zustand vorliegen. Da dort, z.B. durch Gezeitenkräfte, Wärme entsteht, schmelzen diese Stoffe und gelangen zur Oberfläche, wo sie erstarren und sich zu mehreren hundert Meter hohen Ablagerungen aufschichten können.

Methan übernimmt auf Titan eine ähnliche Rolle wie Wasser auf der Erde. Sein Anteil in der Atmosphäre beträgt rund 5 %. Nahe der Oberfläche kann es kondensieren und Wolken bilden, aus denen Niederschlag fällt. Die von der Raumsonde Cassini auf Titan abgesetzte Huygens-Sonde landete im Januar 2005 in einem ausgetrockneten Methan-See. Die Bilder, die sie vorher aus einigen Kilometer Höhe aufgenommen hat, zeigen ein dendritisches Flusssystem, das diesen See einmal gespeist hat (Abb. 2). Auf späteren Radaraufnahmen bei nahen Vorbeiflügen von Cassini an Titan waren zahlreiche spiegelglatte Flächen in polaren Breiten des Mondes zu sehen, die sich kaum anders als als flüssigkeitsgefüllte Senken deuten lassen – eine Seenplatte aus einer Mischung von Methan und anderen Kohlenwasserstoffen.

Aktiven Kryovulkanismus hat Cassini auf dem kleinen Saturnmond Enceladus (Durchmesser 500 km) entdeckt. Veränderungen in der Nähe des Mondes ließen darauf schließen, dass an seinem Südpol riesige Mengen von Staub und Wassereis austreten müssen (Dougherty et al. 2006). Die Instrumente der Sonde zeigten einige tiefe Schluchten, aus denen bis zu 150 kg Wassereismoleküle pro Sekunde austreten

(Waite et al. 2006) – eine Art Kryovulkanismus vergleichbar mit Prozessen in Geysiren oder Schneekanonen auf der Erde (Abb. 3). Damit ließ sich die Quelle des äußersten Rings von Saturn identifizieren und eine der offenen Fragen zum Verständnis des Saturnsystems klären.

Der Dynamo im Planet

Neben der Erde haben mehrere andere Planeten ein eigenes Magnetfeld, das in einem Dynamoprozess entsteht. Grundvoraussetzung dafür ist eine elektrisch leitende und flüssige Region im Inneren des Planeten. Bei den erdähnlichen Planeten handelt es sich um einen Eisenkern, bei Jupiter und Saturn um eine metallische Hochdruckform des Wasserstoffs und bei Uranus und Neptun um eine in der Planetologie als »Eis« bezeichnete Mischung von Wasser, Ammoniak und Methan, die im Inneren der beiden Planeten als ein Fluid mit ionischer Leitfähigkeit vorliegt. Konvektionsströmungen in der leitenden Flüssigkeit verursachen einen Dynamoprozess, bei dem durch Induktion elektrische Ströme und damit ein Magnetfeld entstehen. Für die Erde ist dieser Prozess einigermaßen verstanden, und viele Eigenschaften des Erdmagnetfelds ließen sich in numerischen Modellen reproduzieren (Abb. 4; Christensen und Tilgner 2002).

Betrachten wir die Gesamtheit der Planeten, so erstaunt die Vielfalt in der Stärke und Morphologie ihrer Magnetfelder:

- Jupiters Magnetfeld wird wie das der Erde durch einen leicht gegen die Rotationsachse geneigten Dipolanteil dominiert und ist an der Planetenoberfläche zehnmal stärker als das Erdmagnetfeld.
- Beim Merkur, dem kleinsten und sonnennächsten Planeten, registrierte die Sonde *Mariner 10* in zwei Vorbeiflügen 1974/75 ein Feld, das an der Oberfläche hundertmal schwächer ist als das der Erde. Wegen der geringen Intensität ist zweifelhaft, ob dafür ebenfalls ein Dynamo verantwortlich ist. Im vergangenen Jahr bestätigten Messungen der Raumsonde *Messenger*, dass es sich um ein globales, an der Rotationsachse ausgerichtetes Dipolfeld handelt, was die Dynamohypothese stützt.
- Saturns Dipol scheint völlig parallel zur Rotationsachse ausgerichtet zu sein, auch in den Quadrupol- und Oktupolanteilen konnte die Raumsonde *Cassini*, die seit 2004 den Saturn umkreist, keine Abweichung von der Axialsymmetrie feststellen. Das gibt Rätsel auf,

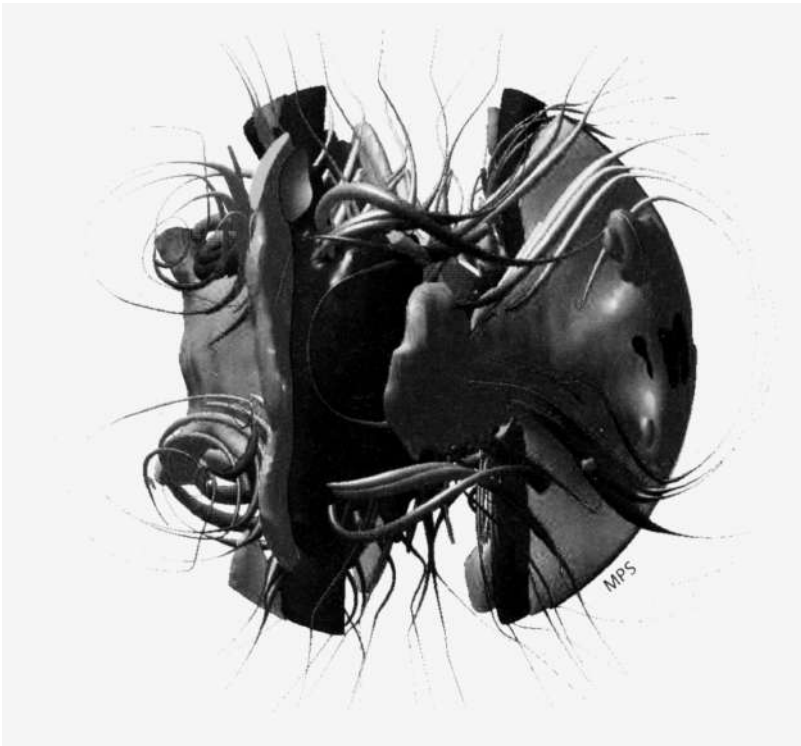


Abbildung 4: Diese Dynamosimulation zeigt die komplexe Struktur magnetischer Feldlinien im flüssigen Eisenkern eines Planeten (Bildnachweis: MPS). (Farbabbildung 9, S. 150)

- da ein grundlegendes Theorem besagt, dass ein homogener Dynamo kein perfekt axialsymmetrisches Magnetfeld erzeugt.
- Uranus und Neptun besitzen Magnetfelder, deren Dipolachsen stark gegen die Rotationsachse geneigt sind und deren Quadrupol- und Oktupolanteile an der Planetenoberfläche genauso stark sind wie der Dipol. Die Satelliten der großen Planeten besitzen kein Magnetfeld, nur der größte Jupitermond Ganymed macht eine Ausnahme.
 - Auch Venus und Mars haben kein eigenes Magnetfeld. 1997 gelangte die Sonde *Mars Global Surveyor* so nahe an die Marsoberfläche heran wie kein Orbiter zuvor. Sie entdeckte starke lokale Magnetfelder in Gebieten, wo die Oberfläche über vier Milliarden Jahre alt ist. Sie lassen sich auf remanente Magnetisierung ferromagnetischer

Mineralien in der Marskruste zurückführen. Diese kann nur in einem starken, von einem Dynamo erzeugten Magnetfeld erworben worden sein, das der Mars in seiner Frühzeit gehabt haben muss.

Aber warum kam der Dynamoprozess zum Erliegen? Die plausibelste Erklärung wäre, dass der Eisenkern des Mars zwar noch flüssig ist, aber nicht konvektiert. Thermische Konvektionsbewegungen setzen voraus, dass die Temperatur nach unten hin stärker ansteigt als mit dem adiabatischen Temperaturgradienten, andernfalls ist die Dichteschichtung stabil. Möglicherweise ist der Wärmefluss aus dem Kern des Mars in der Frühgeschichte des Planeten unter den Wert gesunken, der sich allein durch Wärmeleitung transportieren lässt. Im Gegensatz zum Mars gibt es auf der Erde Plattentektonik, mit der eine höhere Wärmeabgabe aus dem Inneren verbunden ist. Dies könnte der Grund für einen auch heute noch über adiabatischen Temperaturgradienten im Erdkern und die Existenz des Dynamos bei uns sein. Entscheidend könnte aber auch der kleine feste innere Erdkern sein. Selbst wenn die Temperaturschichtung im äußeren flüssigen Kern dynamisch stabil sein sollte, können Konzentrationsunterschiede leichter Elemente eine Art chemische Konvektion antreiben. Der äußere Erdkern besteht zu etwa 10 % aus leichten Komponenten wie Silizium, Sauerstoff oder Schwefel. Im Zuge der Abkühlung der Erde kristallisiert nahezu reines Eisen-Nickel an der Grenze zum inneren Kern aus, und die leichten Komponenten reichern sich in der darüber liegenden Flüssigkeitschicht an. Man nimmt an, dass der Mars (noch) keinen inneren Kern hat, sicher ist dies aber nicht.

Aufschluss hierüber können künftig die Daten von Seismometern liefern, welche die wichtigste Informationsquelle über den inneren Aufbau der Erde sind. Auch die genaue Größe des Marskerns ließe sich so bestimmen. Außer auf der Erde haben Seismometer bisher nur auf dem Mond Daten geliefert. Auf dem Mars könnte ein erstes Instrument mit der GEMS-Mission der NASA zum Einsatz kommen, die sich zur Zeit in einem Auswahlverfahren befindet.

Aktuelle und künftige Raummissionen werden die Magnetfelder einiger Planeten genauer charakterisieren. In einer Verlängerungsphase soll Cassini das Saturnmagnetfeld auch in polaren Breiten unter die Lupe nehmen. Damit lassen sich vielleicht kleine Abweichungen von der Axialsymmetrie finden. Die Juno-Mission der NASA wird die räumliche Struktur des Jupiterfeldes mit einer Genauigkeit bestimmen, die

fast an die heranreicht, mit der das irdische Magnetfeld bekannt ist. Die *Messenger*-Sonde (ab 2011) und die zwei Sonden der europäisch-japanischen *Bepi Colombo Mission* (ab 2019) werden aus Umlaufbahnen um den Merkur dessen Magnetfeld genau vermessen. Erste Dynamo-Modelle für verschiedene Planeten wurden in jüngster Zeit entwickelt (Stanley und Bloxham 2004, Christensen und Wicht 2008). Mit den erwarteten neuen Daten lassen sie sich auf den Prüfstand stellen.

Magnetisches Schutzschild

Die Magnetfelder der Planeten breiten sich im interplanetaren Raum aus und bilden die Magnetosphäre. Diese bildet einen wirksamen Schutzschild gegenüber dem Sonnenwind. Die geladenen Teilchen im interplanetaren Raum werden wirkungsvoll um die Magnetosphäre herumgelenkt und können nicht in die Atmosphäre/Ionosphäre oder auf die Oberfläche gelangen. Nur in den Polregionen des Planeten können sie tiefer eindringen und die wunderschönen Polarlichter auslösen. Wie weit sich die Magnetosphären ausdehnen, hängt von der Stärke des internen Magnetfeldes und von eventuell vorhandenen Teilchenquellen innerhalb der Magnetosphäre ab. In den letzten zehn bis 15 Jahren hat sich das Verständnis der Prozesse in den Magnetosphären von Erde, Jupiter und Saturn dramatisch verbessert.

Vor-Ort-Messungen des Magnetfeldes, der in den Magnetosphären gebundenen geladenen Teilchen und die Messungen elektromagnetischer Wellen haben das Bild der Planetenumgebung stark verändert. Vier Satelliten des Cluster-Projekts messen die Erdmagnetosphäre seit 2000 zeitlich und räumlich aufgelöst. Zwischen 1995 und 2004 flog Galileo als erste Raumsonde in einer Umlaufbahn um Jupiter durch die größte Magnetosphäre unseres Sonnensystems. Dabei konnte sie auch das interne Magnetfeld des Mondes Ganymed und seine Magnetosphäre bestimmen, deren Größe der um Merkur ähnelt. Sie ist in die Magnetosphäre des Jupiter eingebettet und bildet somit einen Sonderfall.

Die Raumsonde Cassini umrundet seit 2004 den Ringplaneten Saturn. Ein wissenschaftliches Ziel dieser Mission ist es, die Saturnmagnetosphäre zu erforschen, die in ihrer Struktur und mit den in ihr ablaufenden physikalischen Prozessen mit der Erde sowie mit Jupiter vergleichbar ist. Neue Ergebnisse von Cassini zeigen, dass in der Magnetosphäre die Zählraten von geladenen und neutralen Teilchen,

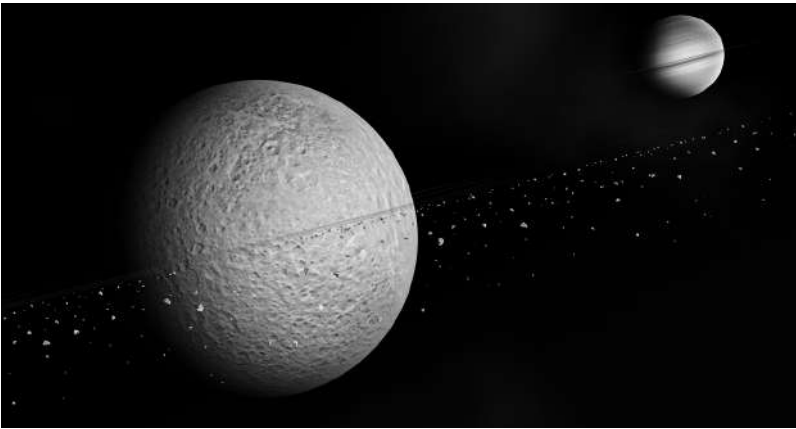


Abbildung 5: Die Ringe um den Saturnmond Rhea (hier als *artist view* gezeigt) wurden durch Absorptionssignaturen in den Zählraten von energiereichen Elektronen entdeckt, die an Bord der Raumsonde Cassini gemessen wurden (Bildnachweis: NASA/JPL, JHUAPL).

das Magnetfeld und die Plasmawellendaten periodisch schwanken. Die Periode dieser Schwankungen entspricht der Rotationsdauer des Planeten und hat vermutlich etwas mit seiner Lage im Sonnensystem zu tun, denn die Rotationsachse und die magnetische Achse sind nahezu parallel zueinander.

Die Galileo-Sonde erforschte Transportphänomene in der Jupitermagnetosphäre und die Wechselwirkung zwischen dem magnetosphärischen Plasma und den Monden im Jupitersystem. Magnetosphärisches Plasma und energiereiche geladene Teilchen bewegen sich entlang und senkrecht zu den Magnetfeldlinien. Kreuzen sich die Bahnen dieser Teilchen irgendwo in der Magnetosphäre mit einem Körper (z.B. ein Mond, Ringe, Gaswolke etc.), so gehen einige oder alle Teilchen verloren. Dadurch entsteht eine Lücke in ihrer Verteilung. Die Auswertung solcher so genannten Absorptionssignaturen ermöglichte es u.a., das Ringsystem um den Saturnmond Rhea zu entdecken (Abb. 5, Jones et al. 2008) und erlaubte den Schluss, dass der Jupitermond Europa einen Wasserstofftorus entlang seiner Bahn um Jupiter ausbildet.

Transportphänomene und Austausch von Materie zwischen verschiedenen Regionen in einer Magnetosphäre sind wichtig, um die

physikalischen Prozesse in Magnetosphären zu verstehen – z.B. die »Injektion« energiereicher (heißer) Teilchen aus den äußeren Regionen in die innere Magnetosphäre mit energieärmerer (kalter) Population. Im Austausch dagegen gelangt Materie über andere Prozesse wieder nach außen. Die Analyse von Teilchenverteilungen in der Jupiter- und Saturnmagnetosphäre aus Daten der Raumsonden Galileo und Cassini hat wesentlich zum Verständnis der ablaufenden Prozesse beigetragen.

Wasser zum Leben

Leben in der Form, wie wir es kennen, ist ohne Wasser undenkbar. Deshalb ist die Suche nach Wasser außerhalb der Erde so wichtig. In den letzten Jahren ließen sich durch Raumsonden und stark verbesserte bodengebundene Beobachtungen einige Kandidaten identifizieren, die heute oder zu früheren Zeiten nachweislich über Wasser verfügen bzw. verfügt haben. Radarmessungen von *Mars Express* zeigen in den Polregionen unter einer Staubschicht Eismassen, die in flüssiger Form den gesamten Planeten mit einem mehrere zehn Meter tiefen Ozean bedecken könnten. Die im Norden des Planeten gelandete Sonde *Phoenix* stieß im letzten Jahr mit einem Baggerarm unter einer sehr dünnen Schicht aus rotem Marsstaub auf Eisbrocken (Abb. 6). Frühere Missionen haben in der Nähe des Marsäquators ausgetrocknete Flusstäler untersucht, in denen vor vielen Millionen Jahren Wasser geflossen sein muss.

Auch die Jupitermonde Europa, Ganymed und Callisto beherbergen vermutlich riesige unterirdische Ozeane aus Wasser, die durch die Wärme aufgrund der starken Gezeitenkräfte zwischen den Monden und Jupiter und durch radioaktive Prozesse im Inneren dieser Körper flüssig bleiben. Ob sich primitive Lebensformen in den Ozeanen bilden konnten, ist bislang unklar. Enceladus mit seinen Geysiren ist ein weiterer Kandidat. Vermutlich sind Wasserreservoirs im Inneren des Mondes für die Fontänen über dem Südpol verantwortlich.

Auch kleinere Körper wie Kometen, transneptunische Objekte oder Asteroiden aus den äußeren Bereichen des Asteroidengürtels beherbergen Wasser und Eis. Möglicherweise haben diese Objekte in der frühen Phase unseres Sonnensystems Wasser auf die Erde und andere Planeten gebracht. Messungen des Mengenverhältnisses von Wasserstoff und seinem Isotop Deuterium deuten eher auf Asteroiden hin, da



Abbildung 6: Mit dem Greifarm grub die *Phoenix*-Sonde am Fuß ihres Landegestells einen hellen, glatten Block aus, den die Wissenschaftler für Eis halten (Bildnachweis: M. Di Lorenzo, K. Kremer, NASA/JPL/UA/MPS/Spaceflight).

in Wassereinschlüssen in kohligen Chondriten (primitiven Meteoriten aus dem Asteroidengürtel) ähnliche Verhältnisse gefunden wurden wie in ozeanischem Wasser. Bisherige Messungen dieses Isotopenverhältnisses an Kometen und transneptunischen Objekten stimmten dagegen schlecht mit irdischem Wasser überein. Unklar ist jedoch, ob die untersuchten Kometen repräsentativ für Kometen aus dem Kuiper-Gürtel sind. Zudem ist inzwischen eine neue Klasse von Kometen im Hauptasteroidengürtel entdeckt worden, die als Quelle für das irdische Wasser infrage kommen (Hsieh et al. 2006).

Für Leben sind auch organische Moleküle erforderlich, die man in Kometenschweifen und in der Atmosphäre des Saturnmondes Titan in großer Anzahl gefunden hat. Mars und die Monde Europa, Enceladus und Titan sind die Körper in unserem Sonnensystem, die in dieser Hinsicht am interessantesten erscheinen. Allerdings beeinträchtigen Umgebungsparameter, wie z.B. die lebensfeindliche hochenergetische Strahlung die Entstehung von Leben im Weltraum wesentlich. Der

Mond Europa liegt zwar in der Magnetosphäre von Jupiter und wird damit von der kosmischen Strahlung und dem Sonnenwind abgeschirmt, aber er umkreist den Planeten in einem Abstand, in dem Elektronen nahezu mit Lichtgeschwindigkeit auf seine Oberfläche treffen und Leben dort unmöglich machen. Selbst gut abgeschirmte elektronische Bauteile von Raumsonden überleben nur wenige Wochen in diesem Teil der Jupitermagnetosphäre – dem stärksten Strahlungsgürtel des Sonnensystems. Mögliches Leben könnte sich hier nur in der Tiefe verbergen.

Weitere Raumsonden werden zu Planeten in unserem Sonnensystem starten und neue Welten entdecken. Dabei bilden Mars und der Erdmond den Schwerpunkt, doch auch Zwergplaneten wie Vesta und Ceres im Asteroidengürtel sind Ziel von Missionen (*Dawn*-Mission). Die Rosetta-Mission der ESA wird nicht nur erstmals einen Kometenkern umkreisen und aus wenigen Kilometer Höhe über viele Monate untersuchen, sondern auch die Landesonde Philae an seiner Oberfläche absetzen.

Literatur

- Christensen U., Tilgner A. (2002) Der Geodynamo. *Physik Journal* 1, 10, S. 41.
- Christensen U., Wicht J. (2008) Models of magnetic field generation in partly stable planetary cores: Applications to Mercury and Saturn. *Icarus* 196, S. 16.
- Dougherty M.K., Khurana K.K., Neubauer F.M., Russell C.T., Saur J., Leisner J.S., Burton M. E. (2006) Identification of a Dynamic Atmosphere at Enceladus with the Cassini Magnetometer. *Science* 311, 5766, S. 1406.
- Harder H., Christensen U. (1996) A one-plume model of martian mantle convection. *Nature* 380, S. 507.
- Neukum G., Jaumann R., Hoffmann H., Hauber E., Head J.W., Basilevsky A.T., Ivanov B.A., Werner S.C., van Gasselt S., Murray J.B., McCord T. & The HRSC Co-Investigator Team (2004) Recent and episodic volcanic and glacial activity on Mars revealed by the High Resolution Stereo Camera. *Nature* 432, S. 971.
- Hsieh H.H., Jewitt D. (2006) A Population of Comets in the Main Asteroid Belt. *Science* 312, 5773, S. 561.
- Jones G.H., Roussos E., Krupp N., Beckmann U., Coates A.J., Cray F., Dandouras I., Dikarev V., Dougherty M.K., Garnier P., Hansen C.J., Hendrix A.R., Hospodarsky G.B., Johnson R.E., Kempf S., Khurana K.K., Krimigis S.M., Krüger H., Kurth W.S., Lagg A., McAndrews H.J., Mitchell D.G., Paranicas C., Postberg F., Russell C.T., Saur J., Seiß M., Spahn F., Srama R., Strobel D.F., Tokar R., Wahlund J.E., Wilson R.J., Woch J., Young D. (2008) The Dust Halo of Saturn's Largest Icy Moon, Rhea. *Science* 319, 5868, S. 1380.
- Stanley S., Bloxham J. (2004) Convective-region geometry as the cause of Uranus' and Neptune's unusual magnetic fields, *Nature* 428, S. 151.
- Waite J.H., Combi M.R., Ip W.H., Cravens T.E., McNutt R.L., Kasprzak W., Yelle R., Luhmann J., Niemann H., Gell D., Magee B., Fletcher G., Lunine J., Tseng W.L. (2006), Cassini Ion and Neutral Mass Spectrometer: Enceladus Plume Composition and Structure, *Science* 311, 5766, S. 1419.

Nina Schaller

Auf Zehenspitzen zum Weltrekord¹

Der Strauß im Fokus moderner Biomechanik

Große Augen, federnder Gang, nun wird keck der Flügel geschüttelt. Imposant steht er vor mir, der größte aller Vögel. Doch da schnellst aus zweieinhalb Metern Höhe sein Schnabel herab, packt meine Mütze und flugs macht sich der Dieb davon. Ich spurte hinterher – ein sinnloses Unterfangen. Zwar sind auch wir Menschen gut zu Fuß: Olympioniken bewältigen in eineinhalb Stunden eine Strecke von 30 Kilometern. Doch mit konstantem Tempo 60 schafft der Strauß das in einem Drittel der Zeit. Mit Spitzengeschwindigkeiten bis 70 Stundenkilometer holt dieser bis zu 150 Kilogramm schwere Vogel fast die schnellsten seiner fliegenden Verwandten ein. Auch Bernhard Grzimek wusste: »Kein anderes Tier läuft so ausdauernd wie der Strauß«. So legt er große Distanzen zurück, um rasch neue Nahrungsgründe zu erreichen oder hungrige Hyänen abzuhängen. Wie schafft es der Strauß, allen anderen davonzulaufen?

Ein ausdauernder Läufer muss wie ein sparsames Auto Antriebsenergie möglichst verlustfrei auf einen effizienten Bewegungsapparat übertragen. Die Antriebsenergie wird dabei durch Stoffwechselvorgänge aus Nahrung und Sauerstoff gewonnen. Das nervengesteuerte Muskel-Skelett-System der Beine setzt diese Energie dann gezielt in Fortbewegung um. Messungen des Sauerstoffverbrauchs haben ergeben, dass rennende Strauße nur zwei Drittel der Energie benötigen, die Vögel solcher Größe rein rechnerisch verbrauchen müssten. Wie der Fortbewegungsapparat diese Energie so sparsam umsetzt, war weitgehend unbekannt. Zwar existierten anatomische Beschreibungen sowie Laborbeobachtungen von Straußen auf Laufbändern. Man kannte also Einzelteile, nicht aber deren Funktion im Gesamtsystem. Mein Ziel war es, Form und Funktion gleichermaßen zu erforschen. Ich suchte

¹ Für diesen Beitrag wurde Nina Schaller 2009 mit dem Klaus Tschira Preis für verständliche Wissenschaft ausgezeichnet (www.klaus-tschira-preis.info).

Beinstrukturen, die Fortbewegung effizienter machen. Um mögliche Auswirkungen auf die Bewegungsdynamik zu erkennen, untersuchte ich parallel dazu lebende Strauße. Während der Ingenieur eine Maschine plant und dann zusammensetzt, musste ich den Strauß »demonstrieren«, um seine dynamischen Prozesse zu verstehen. Erst dann konnte ich die Teile in einen funktionellen Kontext stellen. Die biologischen Methoden hierzu liefern die klassische Morphologie (die Gestaltlehre) und die moderne Biomechanik (die Fortbewegungsanalyse lebender Systeme). Um natürliche Bewegungsabläufe untersuchen zu können, zog ich in einem großen Freigehege drei Strauße auf und gewöhnte sie langsam an mich und die dort installierten Messinstrumente. Für mein fächerübergreifendes Projekt kooperierte ich mit renommierten Morphologen des Forschungsinstituts Senckenberg, der Universität Heidelberg und Biomechanikexperten der Universität Antwerpen.

Der Knick im Bein

Menschen sind Sohlengänger, Vögel gehen nur auf ihren Zehen. Dieser Unterschied ist wichtig, um die Bewegungsdynamik des Straußes zu verstehen (Abb. 1). Der Mittelfuß ist stark verlängert und bildet analog zu unserem Schienbein den unteren Teil des Beins. Dadurch verschiebt sich das Fersengelenk nach oben auf Höhe unseres Knies, während der Unterschenkel auf Position des menschlichen Oberschenkels liegt. Das »Schwungbein« besteht also aus Unterschenkel, Fersengelenk und Mittelfußknochen, weshalb es bei schreitenden Vögeln scheint, als ob das »Knie« nach hinten einknickt.

Das eigentliche Kniegelenk ist angewinkelt und wie der kurze, waagrechte Oberschenkel unter dem Federkleid verborgen. Zunächst habe ich die Beinsegmente verschiedener ans Laufen angepasster Vogelarten, vom hühnergroßen Roadrunner über den 50 Kilogramm schweren Emu bis zum Strauß, vermessen und verglichen. Unter allen gefiederten Laufspezialisten besitzt der Strauß die relativ längsten Beine. Beim Rennen erreicht er Schrittweiten von bis zu vier Metern! Zudem befindet sich die Masse seiner Beinmuskulatur, mehr als bei den anderen Arten, weit oben an Becken und Oberschenkel, während das schwingende Bein am leichtesten ist und über lange Sehnen bewegt wird. Die Grundvoraussetzungen für hohe Laufgeschwindigkeit, große Schrittweite und -frequenz, hat der Strauß damit als Klassenbester

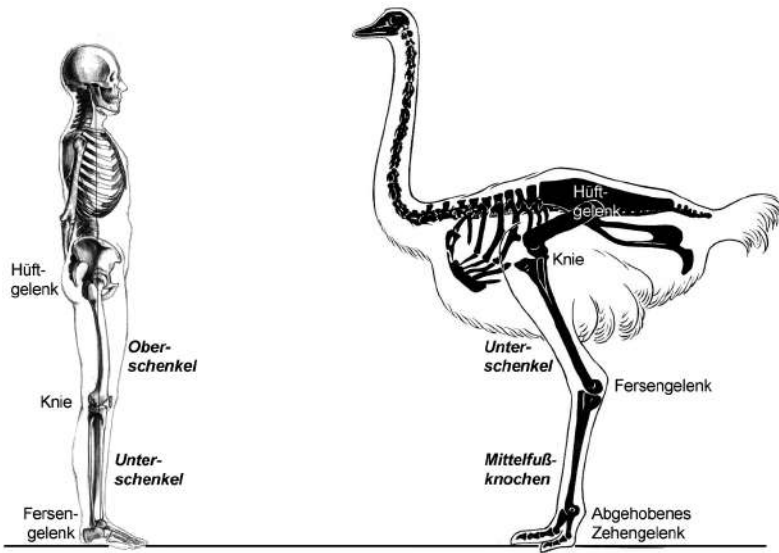


Abbildung 1: Das Hüftgelenk der Vögel, speziell des Straußes, ist im Gegensatz zu unserem Hüftgelenk nur eingeschränkt beweglich. Das Kniegelenk hat wie unser Hüftgelenk einen größeren Bewegungsspielraum, um Veränderungen des Körperschwerpunkts, z.B. bei Richtungswechseln oder beim Abbremsen, auszugleichen. Das Fersengelenk ist wie unser Knie während des Stehens nahezu gestreckt. Dadurch bildet das Bein eine gerade Säule, die das auf ihm lastende Körpergewicht optimal trägt. Während diese Säule beim Menschen aus Ober- und Unterschenkel besteht, wird sie beim Strauß durch Unterschenkel und Mittelfußknochen gebildet.

erfüllt. Je länger das Pendel einer Wanduhr ist, desto weiter schwingt es. Schiebt man das Gewicht dann noch nah an den Drehpunkt, in diesem Fall das Hüftgelenk, erhöht sich auch die Frequenz, mit der das Pendel schwingt.

Ausdauernde Fortbewegung ist auch gekoppelt an verlustfreie Umsetzung von Muskelkraft. Hierzu sollte die Kraft vorrangig in den Vortrieb geleitet und der seitliche Bewegungsspielraum der Beine unterdrückt werden. Gelenkigkeit ermöglicht es uns zwar, auf Bäume zu klettern, beim Geradeausrennen muss ein Mensch diese nun unnötigen Freiheitsgrade aktiv durch Muskelkraft unterdrücken. Bänder können passiv, also energetisch kostenneutral, Gelenkbewegungen limitieren, ähnlich einem Stützkorsett. Um den Anteil der

Bänder an der Bewegungssteuerung zu ermitteln, filmte ich meine Strauße beim Laufen. Danach simulierte und filmte ich im Labor Beinbewegungen mit seziierten Gliedmaßen. Muskeln und Sehnen hatte ich entfernt, so dass nur Bänder die Gelenke zusammenhielten und als mögliche Richtungsgeber übrig waren. Die Auswertung der Filmsequenzen zeigte, dass der seitliche Bewegungsspielraum bei seziierten Beinen fast identisch mit dem lebender Strauße war. Es waren also vorrangig Bänder, die das Bein auf seinem steten Pendelkurs hielten. Muskelkraft wird somit auf Vorschnellen und Abstoßen des Beins konzentriert. Die besondere Anordnung von Seitenbändern erzielte aber noch einen bisher völlig unbekannten Effekt: Bewege ich den langen Mittelfußknochen, um das Fersengelenk zu beugen, musste ich einen Widerstand überwinden – sehr unerwartet bei einem leblosen Bein ohne aktive Muskulatur. Ließ ich den Knochen los, schnappte das Gelenk sogar automatisch in die Strecklage zurück. Beim lebenden Strauß ist das Fersengelenk, wie unser Knie, dann gestreckt, wenn die Zehen am Boden sind, das Bein also das Körpergewicht tragen muss. Für den schweren Strauß bedeutet das einen ständigen Kraftaufwand. Hielten Bänder das Bein passiv aufrecht, würde dies Muskelkraft sparen. Um zu ermitteln, was die Bänder zur Gelenkstabilisierung beitragen, übte ich so lange Druck auf das stehende, seziierte Bein aus, bis das Fersengelenk einknickte. Dazu waren 15 Kilogramm nötig! (Abb. 2) Steht der Strauß auf beiden Beinen, entlasten Bänder die Beinmuskulatur um etwa 20 Prozent. Damit enthält dieser Mechanismus zwei entscheidende Strategien, die die Energieeinsparung bei der Fortbewegung erklären: Er reduziert Muskelkraft und in der Konsequenz Muskelmasse, was wiederum die Pendeleigenschaften des Beins verbessert.

Um auf Dauer schnell zu sein, ist es wichtig, Energieverluste auch am Boden gering zu halten. Dies wird durch Verringerung der Reibungsfläche erreicht. Darum sind die Reifen eines Rennrads wesentlich schmaler als die von Omas Drahtesel. Ein laufstarkes Tier verringert seine Reibungsfläche durch Zehenhaltung und -reduzierung. Das Pferd ist so weit gegangen, dass es nur noch auf dem Nagel seines Mittelfingers, dem Huf, galoppiert. Beim Strauß ist das ähnlich. Während die mit ihm verwandten Laufvögel, Emu und Nandu, drei und alle anderen Vögel vier Zehen nebst Krallen besitzen, hat der Strauß seine Zehenzahl auf zwei reduziert und nur eine Krallen behaltem. Zudem ist er der einzige Vogel, der auf Zehenspitzen geht (Abb. 3).

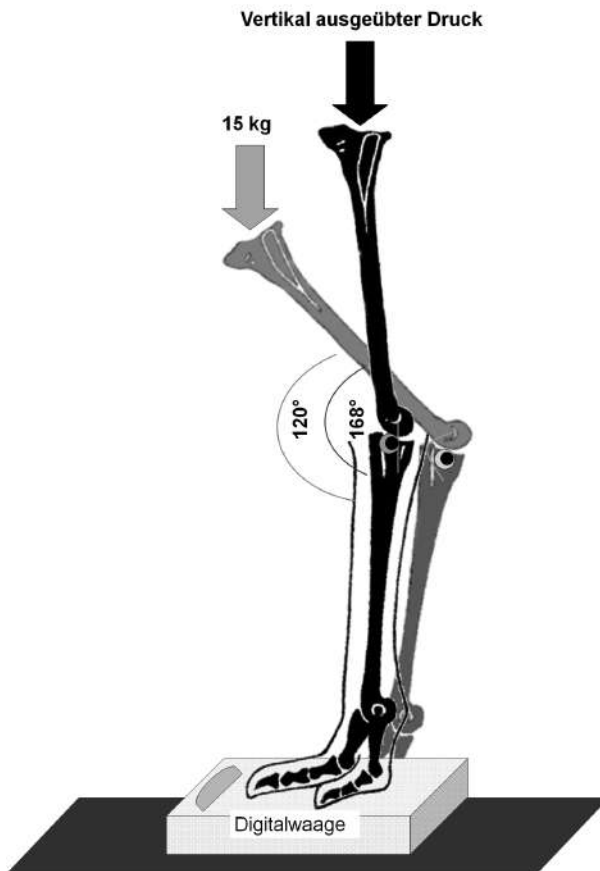


Abbildung 2: Das durch den Einrastmechanismus stabilisierte sezierte Straußenbein kann ein Gewicht von bis zu 15 kg tragen. Dazu tragen der Knochenzapfen und das Seitenband des Fersengelenks bei.

Wenn das Bein vor dem Aufsetzen der Zehen gestreckt wird, muss das Seitenband bei einem Beugewinkel von $\sim 120^\circ$ den nach außen gerichteten Knochenzapfen passieren. Bei diesem kritischen Winkel ist das Band, ähnlich einem Gummiband, äußerst gespannt und rutscht schnell über diesen Zapfen, wodurch das Bein (Unterschenkel und Mittelfußknochen) in die Strecklage schnalzt. Der Knochenzapfen funktioniert während der Standbeinphase wie ein Sperrbolzen, der das Band arretiert und somit das Bein passiv gestreckt hält. Nimmt der Strauß zu Beginn der Schwungbeinphase die Zehen vom Boden und der Beugewinkel des Fersengelenks nähert sich 120° , rutscht das sich entspannende Seitenband ganz leicht über den Knochenzapfen und der untere Teil des Beins (Mittelfußknochen mit Zehen) ist wieder frei beweglich (Farbabbildung 7, S. 148).



Abbildung 3: Vergleich der Füße vom Strauß (a/b), Emu (c) und Reiher (d). Die Zehenzahl beim Strauß ist maximal reduziert: Er verfügt nur über zwei Zehen und eine Krallen, während die mit ihm verwandten Laufvögel Emu und Nandu drei und alle anderen Vögel vier Zehen mit Krallen besitzen. Zudem ist er der einzige Vogel, der auf Zehenspitzen geht. (Bildnachweise: 3 a/b Schaller, 3 c/d Wikimedia Commons).

Wie beim Orthopäden

Die Flächenersparnis ist beachtlich: Der Vergleich von Zehenabdrücken, die ich von Emus, Nandus und Straußen gesammelt hatte, zeigte, dass der Strauß eine 60 Prozent kleinere Zehenfläche hat. Somit sind die Energieverluste durch Reibung beim Strauß extrem gering. Nun wollte ich herausfinden, wie die Zehen mit dem Boden interagieren. Dies hatte man bei lebenden Vögeln noch nie untersucht, eine etablierte Methode existierte also nicht. Ich ließ daher die Strauße über eine Druckmessplatte laufen, wie sie Orthopäden zur Analyse der Druckverteilung menschlicher Füße benutzen. Wie Stoßdämpfer federn die weichen Sohlen der Zehen auch hohe Drucke gleichmäßig ab. Dabei agiert die nach außen gerichtete Zehe wie eine Seitenstütze. Die Krallen, die beim Gehen kaum den Boden berührte, übte beim Rennen Drucke bis zu 40 Kilogramm pro Quadratzentimeter aus (Abb. 4). Dadurch bohrt sie sich wie ein Spike in den Boden und gewährleistet,

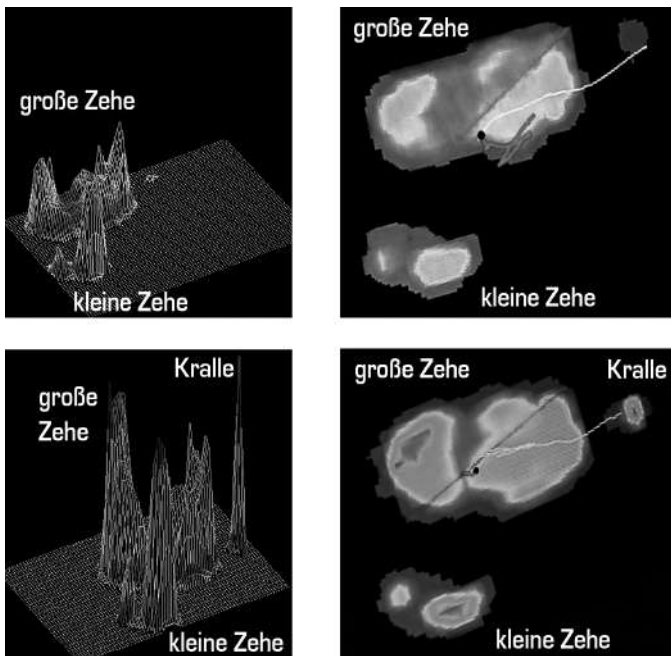


Abbildung 4: Druckverteilung der beiden Straußenzehen (rechter Fuß) beim Gehen und Rennen in drei- und zweidimensionaler Ansicht. Beim Gehen (obere Reihe) wird die Krallen der großen Zehe kaum belastet, während beim Rennen (untere Reihe) sehr hohe Drucke auf ihr lasten: Die Krallen funktioniert bei schnellem Rennen wie ein Positionsanker. Typisch für alle Zweibeiner ist, dass beim Rennen höhere Drucke auftreten als beim Gehen. Wenn Menschen und Vögel gehen, ist immer ein Fuß für mehr als die Hälfte der Schrittdauer am Boden. Es sind also zu einem gegebenen Zeitpunkt immer zwei Füße am Boden, die das Körpergewicht tragen. Beim Rennen wird der Körper durch kräftigeres Abstoßen beschleunigt und die daraus resultierende Flugphase vergrößert die Schrittlänge. Es ist also immer nur ein Fuß am Boden, auf dem das gesamte Körpergewicht lastet (Bildnachweis: Schaller, Farbabbildung 8, S. 149).

dass der Strauß auch bei 70 Stundenkilometern nicht die Bodenhaftung verliert. Der Gang auf Zehenspitzen und die besondere Zehenarchitektur erlauben ein Höchstmaß an Funktionalität bei minimalem Material- und Kosteneinsatz, ideal für ausdauerndes Laufen auf ebenem Savannenboden.



Abbildung 5: Ein Strauß im Neuschnee beim abrupten Abbremsen. Die gleichzeitig ausgebreiteten Flügel funktionieren wie ein Bremsfallschirm (Bildnachweis: Schaller).

Vorläufiges Fazit

Extremleistungen in der Tierwelt haben Menschen seit jeher fasziniert und zum Nachahmen animiert. Moderne Biomechanik und traditionelle morphologische Methoden erlauben neue Einblicke in den Bauplan des schnellsten aller Marathonläufer, des Afrikanischen Straußes.

Der Frage, wie es der Strauß schafft, so lange und so schnell zu rennen, sind wir einen Riesenschritt näher gekommen. Mit optimierten Pendeleigenschaften, passiver Gelenkstabilisierung und minimaler Reibungsfläche funktioniert sein Bewegungsapparat gemäß Energieeffizienzklasse A. Da diese Strategien mechanisch basiert sind, könnten sie vom Menschen nachgebaut und technisch genutzt werden. Es gibt bereits Ansätze, diese Erkenntnisse in der Prothetik und für den Bau zweibeiniger Roboter anzuwenden.

Beschwingt in die Zukunft

Der Strauß lehrt aber auch, dass die Natur nicht verschwenderisch ist. In den vielen Jahren unserer Zusammenarbeit beobachtete ich, dass meine gefiederten Arbeitsgruppenmitglieder ihre großen Flügel bei rasanten Wendemanövern als Seitenruder benutzen (Abb. 5). Da Strauße flugunfähig sind, nahmen Forscher bisher an, dass die Flügel Rudimente und in der Rückbildung begriffen seien. Laut meinen Beobachtungen sowie mit Kooperationspartnern durchgeführten Windkanalexperimenten und Modellrechnungen sind sie für den Strauß jedoch unentbehrlich. Seine langen Beine haben sich in 60 Millionen Jahren auf schnelles Geradeauslaufen spezialisiert. Für flinke Zickzack-Manöver sind die »Hochleistungsstelzen« jedoch nicht gelenkig genug – diesen Nachteil gleichen die Flügel als Balancierorgane aus. Meine neuesten Erkenntnisse sind auch für Paläontologen interessant: Auch die Vorfahren heutiger Vögel, zweibeinige Dinosaurier, könnten ihre befiederten Vorderarme zur Vergrößerung ihres Bewegungsspektrums eingesetzt haben und nicht nur, wie bisher angenommen, zum Fangen von Insekten oder für Imponiergehabe. Möglicherweise kann uns gerade ein flugunfähiger Vogel einiges über die Evolution des Vogelfluges lehren. Für mich und meine Wissenschaftsstrauße gibt es also noch viel zu entdecken.

Literatur

Schaller N.U., Herkner B., Villa R., Aerts P. (2009) The intertarsal joint of the ostrich (*Struthio camelus*). Anatomical examination and function of passive structures in locomotion. *Journal of Anatomy* 214, S. 830.

Aaron Ciechanover

**Intracellular Protein Degradation:
From a Vague Idea thru the Lysosome
and the Ubiquitin-Proteasome System and
onto Human Diseases and Drug Targeting**

(Nobel-Vortrag 2004)¹

Between the 1950s and 1980s, scientists were focusing mostly on how the genetic code was transcribed to RNA and translated to proteins, but how proteins were degraded had remained a neglected research area. With the discovery of the lysosome by Christian de Duve it was assumed that cellular proteins are degraded within this organelle. Yet, several independent lines of experimental evidence strongly suggested that intracellular proteolysis was largely non-lysosomal, but the mechanisms involved have remained obscure. The discovery of the ubiquitin-proteasome system resolved the enigma. We now recognize that degradation of intracellular proteins is involved in regulation of a broad array of cellular processes, such as cell cycle and division, regulation of transcription factors, and assurance of the cellular quality control. Not surprisingly, aberrations in the system have been implicated in the pathogenesis of human disease, such as malignancies and neurodegenerative disorders, which led subsequently to an increasing effort to develop mechanism-based drugs.

- 1 Wir danken der Nobel-Stiftung für die Genehmigung des Wiederabdrucks aus Les Prix Nobel. The Nobel Prizes 2004, hrsg. v. Tore Frängsmyr, Stockholm 2005, S. 151 ff.

Abbreviations used: ODC, ornithine decarboxylase; G6PD, glucose-6-phosphate dehydrogenase; PEPCK, phosphoenol-pyruvate carboxykinase; TAT, tyrosine aminotransferase; APF-1, ATP-dependent Proteolysis Factor 1 (ubiquitin); UBIP, ubiquitous immunopoietic polypeptide (ubiquitin); MCP, multicatalytic proteinase complex (26S proteasome); CP, 20S core particle (of the 26S proteasome); RP, 19S regulatory particle (of the 26S proteasome).

Introduction

The concept of protein turnover is hardly 70 years old. Beforehand, body proteins were viewed as essentially stable constituents that were subject to only minor »wear and tear«: dietary proteins were believed to function primarily as energy-providing fuel, which were independent from the structural and functional proteins of the body. The problem was hard to approach experimentally, as research tools were not available. Important research tools that were lacking at that time were stable isotopes. While radioactive isotopes were developed earlier by George de Hevesy (de Hevesy G., *Chemistry* 1943. In: *Nobel Lectures in Chemistry 1942-1962*. World Scientific 1999. pp. 5-41), they were mostly unstable and could not be used to follow metabolic pathways). The concept that body structural proteins are static and the dietary proteins are used only as a fuel was challenged by Rudolf Schoenheimer in Columbia University in New York City. Schoenheimer escaped from Germany and joined the Department of Biochemistry in Columbia University founded by Hans T. Clarke (1-3). There he met Harold Urey who was working in the Department of Chemistry and who discovered deuterium, the heavy isotope of hydrogen, a discovery that enabled him to prepare heavy water, D₂O. David Rittenberg who had recently received his Ph.D. in Urey's laboratory, joined Schoenheimer, and together they entertained the idea of »employing a stable isotope as a label in organic compounds, destined for experiments in intermediary metabolism, which should be biochemically indistinguishable from their natural analog« (1). Urey later succeeded in enriching nitrogen with ¹⁵N, which provided Schoenheimer and Rittenberg with a »tag« for amino acids and as a result for the study of protein dynamics. They discovered that following administration of ¹⁵N-labeled tyrosine to rat, only ~50 % can be recovered in the urine, »while most of the remainder is deposited in tissue proteins. An equivalent of protein nitrogen is excreted« (4). They further discovered that from the half that was incorporated into body proteins »only a fraction was attached to the original carbon chain, namely to tyrosine, while the bulk was distributed over other nitrogenous groups of the proteins« (4), mostly as an αNH₂ group in other amino acids. These experiments demonstrated unequivocally that the body structural proteins are in a dynamic state of synthesis and degradation, and that even individual amino acids are in a state of dynamic interconver-

sion. Similar results were obtained using ^{15}N -labeled leucine (5). This series of findings shattered the paradigm in the field at that time that: (1) ingested proteins are completely metabolized and the products are excreted, and (2) that body structural proteins are stable and static. Schoenheimer was invited to deliver the prestigious Edward K. Dunham lecture at Harvard University where he presented his revolutionary findings. After his untimely tragic death in 1941, his lecture notes were edited Hans Clarke, David Rittenberg and Sarah Ratner, and were published in a small book by Harvard University Press. The editors called the book »The Dynamic State of Body Constituents«, adopting the title of Schoenheimer's presentation. In the book, the new hypothesis was clearly presented:

»The simile of the combustion engine pictured the steady state flow of fuel into a fixed system, and the conversion of this fuel into waste products. The new results imply that not only the fuel, but the structural materials are in a steady state of flux. The classical picture must thus be replaced by one which takes account of the dynamic state of body structure« (6).

However, the idea that proteins are turning over had not been accepted easily, and was challenged as late as the mid-1950s. For example, Hogness and colleagues studied the kinetics of β -galactosidase in *E. coli* and summarized their findings:

»To sum up: there seems to be no conclusive evidence that the protein molecules within the cells of mammalian tissues are in a dynamic state. Moreover, our experiments have shown that the proteins of growing *E. coli* are static. Therefore it seems necessary to conclude that the synthesis and maintenance of proteins within growing cells is not necessarily or inherently associated with a »dynamic state« (7).

While the experimental study involved the bacterial β -galactosidase, the conclusions were broader, including also the authors' hypothesis on mammalian proteins. The use of the term »dynamic state« was not incidental, as they challenged directly Schoenheimer's studies.

Now, after more than seven decades of research in the field of intracellular proteolysis, and with the discovery of the lysosome and later the ubiquitin-proteasome system, it is clear that the field has been revolutionized. We now recognize that intracellular proteins are turn-

ing over extensively, that the process is specific, and that the stability of many proteins is regulated individually and can vary under different conditions. From a scavenger, unregulated and non-specific end process, it has become clear that proteolysis of cellular proteins is a highly complex, temporally controlled and tightly regulated process that plays major roles in a broad array of basic pathways. Among these processes are cell cycle, development, differentiation, regulation of transcription, antigen presentation, signal transduction, receptor-mediated endocytosis, quality control, and modulation of diverse metabolic pathways. Subsequently, it has changed the paradigm that regulation of cellular processes occurs mostly at the transcriptional and translational levels, and has set regulated protein degradation in an equally important position. With the multitude of substrates targeted and processes involved, it has not been surprising to find that aberrations in the pathway have been implicated in the pathogenesis of many diseases, among them certain malignancies, neurodegeneration, and disorders of the immune and inflammatory system. As a result, the system has become a platform for drug targeting, and mechanism-based drugs are currently developed, one of them is already on the market.

The lysosome and intracellular protein degradation

In the mid-1950s, Christian de Duve discovered the lysosome (see, for example, Refs. 8 and 9 and Figure 1). The lysosome was first recognized biochemically in rat liver as a vacuolar structure that contains various hydrolytic enzymes which function optimally at an acidic pH. It is surrounded by a membrane that endows the contained enzymes latency that is required to protect the cellular contents from their action (see below). The definition of the lysosome was broadened over the years because it had been recognized that the digestive process is dynamic and involves numerous stages of lysosomal maturation together with the digestion of both exogenous proteins (which are targeted to the lysosome through receptor-mediated endocytosis and pinocytosis) and exogenous particles (which are targeted via phagocytosis; the two processes are known as heterophagy), as well as digestion of endogenous proteins and cellular organelles (which are targeted by micro- and macro-autophagy; see Figure 2). The lysosomal/vacu-

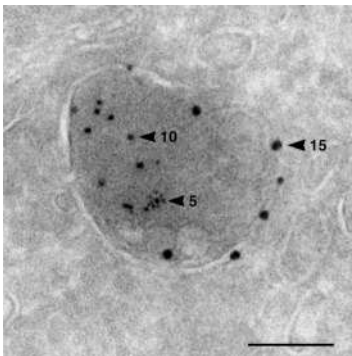


Figure 1: The lysosome. Ultrathin cryosection of a rat PC12 cell that had been loaded for 1 hour with bovine serum albumin (BSA)-gold (5 nm particles) and immunolabeled for the lysosomal enzyme cathepsin B (10-nm particles) and the lysosomal membrane protein LAMP1 (15 nm particles). Lysosomes are recognized also by their typical dense content and multiple internal membranes. Bar, 100 nm (Courtesy of Viola Oorschot and Judith Klumperman, Department of Cell Biology, University Medical Centre Utrecht, The Netherlands).

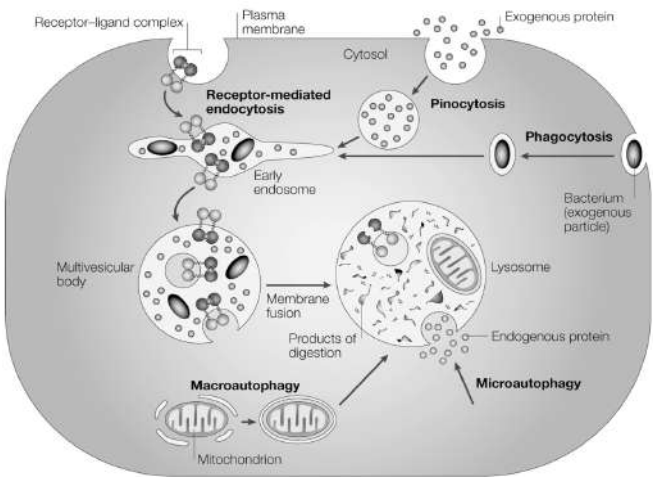


Figure 2: The four digestive processes mediated by the lysosome. (i) specific receptor-mediated endocytosis, (ii) pinocytosis (non-specific engulfment of cytosolic droplets containing extracellular fluid), (iii) phagocytosis (of extracellular particles), and (iv) autophagy (micro- and macro-; of intracellular proteins and organelles) (with permission from Nature Publishing Group. Published originally in Ref. 83).

olar system as we currently recognize it is a discontinuous and heterogeneous digestive system that also includes structures that are devoid of hydrolases – for example, early endosomes which contain endocytosed receptor-ligand complexes and pinocytosed/phagocytosed extracellular contents. On the other extreme it includes the residual bodies – the end products of the completed digestive processes of heterophagy and autophagy. In between these extremes one can observe: primary/nascent lysosomes that have not been engaged yet in any proteolytic process; early autophagic vacuoles that might contain intracellular organelles; intermediate/late endosomes and phagocytic vacuoles (heterophagic vacuoles) that contain extracellular contents/particles; and multivesicular bodies (MVBs) which are the transition vacuoles between endosomes/phagocytic vacuoles and the digestive lysosomes.

The discovery of the lysosome along with independent experiments that were carried out at the same time and that have further strengthened the notion that cellular proteins are indeed in a constant state of synthesis and degradation (see, for example, Ref. 10), led scientists to feel, for the first time, that they have at hand an organelle that can potentially mediate degradation of intracellular proteins. The fact that the proteases were separated from their substrates by a membrane provided an explanation for controlled degradation, and the only problem left to be explained was how the substrates are translocated into the lysosomal lumen, exposed to the activity of the lysosomal proteases and degraded. An important discovery in this respect was the unraveling of the basic mechanism of action of the lysosome-autophagy (reviewed in Ref. 11). Under basal metabolic conditions, portions of the cytoplasm, which contain the entire cohort of cellular proteins, are segregated within a membrane-bound compartment, and are then fused to a primary nascent lysosome and their contents digested. This process was called microautophagy. Under more extreme conditions, starvation for example, mitochondria, endoplasmic reticulum membranes, glycogen bodies and other cytoplasmic entities, can also be engulfed by a process called macroautophagy (see, for example, Ref. 12; the different modes of action of the lysosome in digesting extra- and intracellular proteins are shown in Figure 2).

However, over a period of more than two decades, between the mid-1950s and the late 1970s, it has become gradually more and more difficult to explain several aspects of intracellular protein degradation based on the known mechanisms of lysosomal activity: accumulating

lines of independent experimental evidence indicated that the degradation of at least certain classes of cellular proteins must be non-lysosomal. Yet, in the absence of any »alternative«, researchers came with different explanations, some more substantiated and others less, to defend the »lysosomal« hypothesis.

First was the gradual discovery that came from different laboratories, that different proteins vary in their stabilities, and their half-life times can span three orders of magnitude, from a few minutes to many days. Thus, the $t_{1/2}$ of ornithine decarboxylase (ODC) is ~10 min, while that of glucose-6-phosphate dehydrogenase (G6PD) is 15 hours (for review articles, see, for example, Refs. 13, 14). Also, rates of degradation of many proteins were shown to change with changing physiological conditions, such as availability of nutrients or hormones. It was conceptually difficult to reconcile the findings of distinct and changing half lives of different proteins with the mechanism of action of the lysosome, where the microautophagic vesicle contains the entire cohort of cellular (cytosolic) proteins that are therefore expected to degrade at the same rate. Similarly, changing pathophysiological conditions, such as starvation or re-supplementation of nutrients, were expected to affect the stability of all cellular proteins to the same extent. Clearly, this was not the case.

Another source of concern about the lysosome as the organelle in which intracellular proteins are degraded were the findings that specific and general inhibitors of lysosomal proteases have different effects on different populations of proteins, making it clear that distinct classes of proteins are targeted by different proteolytic machineries. Thus, the degradation of endocytosed/pinocytosed extracellular proteins was significantly inhibited, a partial effect was observed on the degradation of long-lived cellular proteins, and almost no was detected on the degradation of short-lived and abnormal/mutated proteins.

Finally, the thermodynamically paradoxical observation that the degradation of cellular proteins requires metabolic energy, and more importantly, the emerging evidence that the proteolytic machinery uses the energy directly, were in contrast with the known mode of action of lysosomal proteases that under the appropriate acidic conditions, and similar to all known proteases, degrade proteins in an exergonic manner.

The assumption that the degradation of intracellular proteins is mediated by the lysosome was nevertheless logical. Proteolysis results

from direct interaction between the target substrates and proteases, and therefore it was clear that active proteases cannot be free in the cytosol which would have resulted in destruction of the cell. Thus, it was recognized that any suggested proteolytic machinery that mediates degradation of intracellular protein degradation must also be equipped with a mechanism that separates – physically or virtually – between the proteases and their substrates, and enables them to associate only when needed. The lysosomal membrane provided this fencing mechanism. Obviously, nobody could have predicted that a new mode of post-translational modification – ubiquitination – could function as a proteolysis signal, and that untagged proteins will remain protected. Thus, while the structure of the lysosome could explain the separation necessary between the proteases and their substrates, and autophagy could explain the mechanism of entry of cytosolic proteins into the lysosomal lumen, major problems have remained unsolved. Important among them were: (i) the varying half lives, (ii) the energy requirement, and (iii) the distinct response of different populations of proteins to lysosomal inhibitors. Thus, according to one model, it was proposed that different proteins have different sensitivities to lysosomal proteases, and their half lives *in vivo* correlate with their sensitivity to the action of lysosomal proteases *in vitro* (15). To explain an extremely long half-life of a protein that was nevertheless sensitive to lysosomal proteases, or alterations in the stability of a single protein under various physiological states, it was suggested that although all cellular proteins are engulfed into the lysosome, only the short-lived proteins are degraded, whereas the long-lived proteins exit back into the cytosol:

»To account for differences in half-life among cell components or of a single component in various physiological states, it was necessary to include in the model the possibility of an exit of native components back to the extralysosomal compartment« (16).

According to a different model, selectivity was determined by the binding affinity of the different proteins to the lysosomal membrane which controls their entry rates into the lysosome, and subsequently their degradation rates (17). For a selected group of proteins, such as the gluconeogenic enzymes phosphoenol-pyruvate carboxykinase (PEPCK) and fructose-1,6-biphosphatase, it was suggested, though not firmly substantiated, that their degradation in the yeast vacuole was

regulated by glucose via a mechanism called »catabolite inactivation« that possibly involves their phosphorylation. However this regulated mechanism for vacuolar degradation was limited only to a small and specific group of proteins (see for example Ref. 18; reviewed in Ref. 19). More recent studies have shown that at least for stress-induced macroautophagy, a general sequence of amino acids, KFFERQ, directs, via binding to a specific »receptor« and along with cytosolic and lysosomal chaperones, the regulated entry of many cytosolic proteins into the lysosomal lumen. While further corroboration of this hypothesis is still required, it can only explain the mass entry of a large population of proteins that contain a homologous sequence, but not the targeting for degradation of a specific protein under defined conditions (reviewed in refs. 20, 21). The energy requirement for protein degradation was described as indirect, and necessary, for example, for protein transport across the lysosomal membrane (22) and/or for the activity of the H^+ pump and the maintenance of the low acidic intralysosomal pH that is necessary for optimal activity of the proteases (23). We now know that both mechanisms require energy. In the absence of any alternative, and with lysosomal degradation as the most logical explanation for targeting all known classes of proteins at the time, Christian de Duve summarized his view on the subject in a review article published in the mid-1960s, saying: »Just as extracellular digestion is successfully carried out by the concerted action of enzymes with limited individual capacities, so, we believe, is intracellular digestion« (24). The problem of different sensitivities of distinct protein groups to lysosomal inhibitors has remained unsolved, and may have served as an important trigger in future quest for a non-lysosomal proteolytic system.

Progress in identifying the elusive, non-lysosomal proteolytic system(s) was hampered by the lack of a cell-free preparation that could faithfully replicate the cellular proteolytic events – i.e. degrading proteins in a specific and energy-requiring mode. An important breakthrough was made by Rabinovitz and Fisher who found that rabbit reticulocytes degrade abnormal, amino acid analogue-containing hemoglobin (25). Their experiments modeled known disease states, the hemoglobinopathies. In these diseases abnormal mutated hemoglobin chains (such as sickle cell hemoglobin) or excess of unassembled normal hemoglobin chains (which are synthesized normally, but also excessively in thalassemias, diseases in which the pairing chain is not synthesized at all or is mutated and rapidly degraded, and consequent-

ly the bi-heterodimeric hemoglobin complex is not assembled) are rapidly degraded in the reticulocyte (26, 27). Reticulocytes are terminally differentiating red blood cells that do not contain lysosomes. Therefore, it was postulated that the degradation of hemoglobin in these cells was mediated by a non-lysosomal machinery. Etlinger and Goldberg (28) were the first to isolate and characterize a cell-free proteolytic preparation from reticulocytes. The crude extract selectively degraded abnormal hemoglobin, required ATP hydrolysis, and acted optimally at a neutral pH, which further corroborated the assumption that the proteolytic activity was of a non-lysosomal origin. A similar system was isolated and characterized later by Hershko, Ciechanover, and their colleagues (29). Additional studies by this group led subsequently to resolution, characterization, and purification of the major enzymatic components from this extracts and to the discovery of the ubiquitin signalling system (see below).

The lysosome hypothesis is challenged

As mentioned above, the unraveled mechanism(s) of action of the lysosome could explain only partially and at times not satisfactorily, several key emerging characteristics of intracellular protein degradation. Among them were the heterogeneous stability of individual proteins, the effect of nutrients and hormones on their degradation, and the dependence of intracellular proteolysis on metabolic energy. The differential effect of selective inhibitors on the degradation of different classes of cellular proteins (see above but mostly below), could not be explained at all.

The evolution of methods to monitor protein kinetics in cells together with the development of specific and general lysosomal inhibitors has resulted in the identification of different classes of cellular proteins, long- and short-lived, and the discovery of the differential effects of the inhibitors on these groups (see, for example, Refs. 30,31). An elegant experiment in this respect was carried out by Brian Poole and his colleagues in the Rockefeller University. Poole was studying the effect of lysosomotropic agents, weak bases such as ammonium chloride and chloroquine, which accumulate in the lysosome and dissipate its low acidic pH. It was assumed that this mechanism underlies also the anti-malarial activity of chloroquine and similar drugs where

they inhibit the activity of the parasite's lysosome, »paralyzing« its ability to digest the host's hemoglobin during the intra-erythrocytic stage of its life cycle. Poole and his colleagues metabolically labeled endogenous proteins in living macrophages with ^3H -leucine and »fed« them with dead macrophages that had been previously labeled with ^{14}C -leucine. They assumed, apparently correctly, that the dead macrophages debris and proteins will be phagocytosed by live macrophages and targeted to the lysosome for degradation. They monitored the effect of lysosomotropic agents on the degradation of these two protein populations. In particular, they studied the effect of the weak bases chloroquine and ammonium chloride (which enter the lysosome and neutralize the H^+ ions), and the acid ionophore X537A which dissipates the H^+ gradient across the lysosomal membrane. They found that these drugs specifically inhibited the degradation of extracellular proteins, but not that of intracellular proteins (32). Poole summarized these experiments and explicitly predicted the existence of a non-lysosomal proteolytic system that degrades intracellular proteins:

»Some of the macrophages labeled with tritium were permitted to endocytise the dead macrophages labeled with ^{14}C . The cells were then washed and replaced in fresh medium. In this way we were able to measure in the same cells the digestion of macrophage proteins from two sources. The exogenous proteins will be broken down in the lysosomes, while the endogenous proteins will be broken down wherever it is that endogenous proteins are broken down during protein turnover« (33; the paragraph is copied verbatim; A.C.).

The requirement for metabolic energy for the degradation of both prokaryotic (34) and eukaryotic (10, 35) proteins was difficult to explain. Proteolysis is an exergonic process and the thermodynamically paradoxical energy requirement for intracellular proteolysis made researchers believe that energy cannot be consumed directly by proteases or the proteolytic process *per se*, and is used indirectly. As Simpson summarized his findings:

»The data can also be interpreted by postulating that the release of amino acids from protein is itself directly dependent on energy supply. A somewhat similar hypothesis, based on studies on autolysis in tissue minces, has recently been advanced, but the supporting

data are very difficult to interpret. However, the fact that protein hydrolysis as catalyzed by the familiar proteases and peptidases occurs exergonically, together with the consideration that autolysis in excised organs or tissue minces continues for weeks, long after phosphorylation or oxidation ceased, renders improbable the hypothesis of the direct energy dependence of the reactions leading to protein breakdown« (10).

Being cautious however, and probably unsure about this unequivocal conclusion, Simpson still left a narrow orifice opened for a proteolytic process that requires energy in a direct manner: »However, the results do not exclude the existence of two (or more) mechanisms of protein breakdown, one hydrolytic, the other energy-requiring«. Since any proteolytic process must be at one point or another hydrolytic, the statement that makes a distinction between a hydrolytic process, and an energy-requiring, yet non-hydrolytic one, is not clear. Judging the statement from an historical point of view and knowing the mechanism of action of the ubiquitin system, where energy is required also in the pre-hydrolytic step (ubiquitin conjugation), Simpson may have thought of a two step mechanism, but did not give it a clear description. At the end of this clearly understandable, but at the same time difficult and convoluted deliberation, Simpson left us with a vague explanation linking protein degradation to protein synthesis, a process that was known to require metabolic energy:

»The fact that a supply of energy seems to be necessary for both the incorporation and the release of amino acids from protein might well mean that the two processes are interrelated. Additional data suggestive of such a view are available from other types of experiments. Early investigations on nitrogen balance by Benedict, Folin, Gamble, Smith, and others point to the fact that the rate of protein catabolism varies with the dietary protein level. Since the protein level of the diet would be expected to exert a direct influence on synthesis rather than breakdown, the altered catabolic rate could well be caused by a change in the rate of synthesis« (10).

With the discovery of lysosomes in eukaryotic cells it could be argued that energy was required for the transport of substrates into the lysosome or for maintenance of the low intralysosomal pH (see above), for example. The observation by Hershko and Tomkins that the activity

of tyrosine aminotransferase (TAT) was stabilized following depletion of ATP (35) indicated that energy could be required at an early stage of the proteolytic process, most probably before proteolysis occurs. Yet, it did not provide a clue as for the mechanism involved: energy could be used, for example, for specific modification of TAT, e.g. phosphorylation, that would sensitize it to degradation by the lysosome or by a yet unknown proteolytic mechanism, or for a modification that activates its putative protease. It could also be used for a more general lysosomal mechanism, one that involves transport of TAT into the lysosome, for example. The energy inhibitors inhibited almost completely degradation of the entire population of cell proteins, confirming previous studies (e.g. 10) and suggesting a general role for energy in protein catabolism. Yet, an interesting finding was that energy inhibitors had an effect that was distinct from that of protein synthesis inhibitors which affected only enhanced degradation (induced by steroid hormone depletion), but not basal degradation. This finding ruled out, at least partially, a tight linkage between protein synthesis and degradation. In bacteria, which lack lysosomes, an argument involving energy requirement for lysosomal degradation could not have been proposed, but other indirect effects of ATP hydrolysis could have affected proteolysis in *E. coli*, such as phosphorylation of substrates and/or proteolytic enzymes, or maintenance of the »energized membrane state«. According to this model, proteins could become susceptible to proteolysis by changing their conformation, for example, following association with the cell membrane that maintains a local, energy-dependent gradient of a certain ion. While such an effect was ruled out (37), and since there was no evidence for a phosphorylation mechanism (although the proteolytic machinery in prokaryotes had not been identified at that time), it seemed that at least in bacteria, energy was required directly for the proteolytic process. In any event, the requirement for metabolic energy for protein degradation in both prokaryotes and eukaryotes, a process that is exergonic thermodynamically, strongly indicated that in cells proteolysis is highly regulated, and that a similar principle/mechanism has been preserved along evolution of the two kingdoms. Implying from the possible direct requirement for ATP in degradation of proteins in bacteria, it was not too unlikely to assume a similar direct mechanism in the degradation of cellular proteins in eukaryotes. Supporting this notion was the description of the cell-free proteolytic system in reticulocytes (28, 29), a cell that lacks

lysosomes, which indicates that energy is probably required directly for the proteolytic process, although here too, the underlying mechanisms had remained enigmatic at the time. Yet, the description of the cell-free system paved the road for detailed dissection of the underlying mechanisms involved.

The ubiquitin-proteasome system

The cell-free proteolytic system from reticulocytes (28, 29) turned out to be an important and rich source for the purification and characterization of the enzymes that are involved in the ubiquitin-proteasome system. Initial fractionation of the crude reticulocyte cell extract on the anion-exchange resin diethylaminoethyl cellulose yielded two fractions which were both required to reconstitute the energy-dependent proteolytic activity that is found in the crude extract: The unadsorbed, flow through material was denoted fraction I, and the high salt eluate of the adsorbed proteins which was denoted fraction II (see Ref. 38, and Table 1).

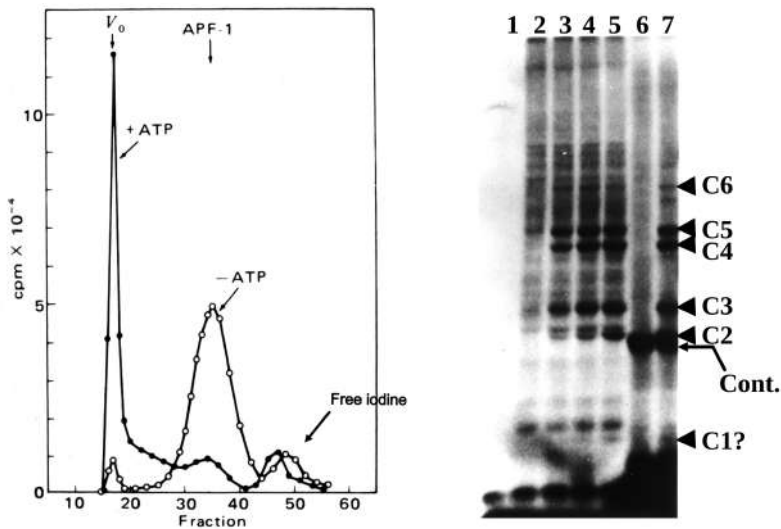
Fraction	Degradation of [³ H]globin (%)	
	– ATP	+ ATP
Lysate	1.5	10
Fraction I	0.0	0.0
Fraction II	1.5	2.7
Fraction I and Fraction II	1.6	10.6

Table 1: Resolution of the ATP-dependent proteolytic activity from crude reticulocyte extract into two essentially required complementing activities (adapted from Ref. 38; with permission from Elsevier/Biochem. Biophys. Res. Commun.).

This was an important observation and a lesson for the future dissection of the system. For one it suggested that the system was not composed of a single ›classical‹ protease that has evolved evolutionarily to acquire energy dependence (although such energy-dependent proteases, the mammalian 26S proteasome (see below) and the prokaryotic

Lon gene product have been described later), but that it was made of at least two components. This finding of a two component, energy-dependent protease, left the researchers with no paradigm to follow, and in attempts to explain the finding, they suggested, for example, that the two fractions could represent an inhibited protease and its activator. Second, learning from this reconstitution experiment and the essential dependence between the two active components, we continued to reconstitute activity from resolved fractions whenever we encountered a loss of activity along further purification steps. This biochemical »complementation« approach resulted in the discovery of additional enzymes of the system, all required to be present in the reaction mixture in order to catalyze the multi-step proteolysis of the target substrate. We chose first to purify the active component from fraction I. It was found to be a small, ~8.5 kDa heat stable protein that was designated ATP-dependent Proteolysis Factor 1, APF-1. APF-1 was later identified as ubiquitin (see below; I am using the term APF-1 to the point where it was identified as ubiquitin and then change terminology accordingly). In retrospect, the decision to start the purification efforts with fraction I turned out to be important, as fraction I contained only one single protein – APF-1 – that was necessary to stimulate proteolysis of the model substrate we used at the time, while fraction II turned out to contain many more. Later studies showed that fraction I contains other components necessary for the degradation of other substrates, but these were not necessary for the reconstitution of the system at that time. This enabled us not only to purify APF-1, but also to quickly decipher its mode of action. If we would have started our purification efforts with fraction II, we would have encountered a significantly bumpier road. A critically important finding that paved the way for future developments in the field was that multiple moieties of APF-1 are covalently conjugated to the target substrate when incubated in the presence of fraction II, and the modification requires ATP (39,40; Figures 3 and 4). It was also found that the modification is reversible, and APF-1 could be removed from the substrate or its degradation products (40).

The discovery that APF-1 was covalently conjugated to protein substrates and stimulates their proteolysis in the presence of ATP and crude fraction II, led in 1980 to the proposal of a model according to which protein substrate modification by multiple moieties of APF-1 targets it for degradation by a downstream, at that time a yet uniden-



Left: Figure 3: APF-1/Ubiquitin is shifted to high molecular mass compound(s) following incubation in ATP-containing crude cell extract. ^{125}I -labelled APF-1/Ubiquitin was incubated with reticulocyte crude Fraction II in the absence (open circles) or presence (closed circles) of ATP, and the reaction mixtures were resolved via gel filtration chromatography. Shown is the radioactivity measured in each fraction. As can be seen, following addition of ATP, APF-1/ubiquitin becomes covalently attached to some component(s) in fraction II, which could be another enzyme of the system or its substrate(s) (with permission from Proceedings of the National Academy of the USA. Published originally in Ref. 39).

Right: Figure 4: Multiple molecules of APF-1/Ubiquitin are conjugated to the proteolytic substrate, probably signalling it for degradation. To interpret the data described in the experiment depicted in Figure 2 and to test the hypothesis that APF-1 is conjugated to the target proteolytic substrate, ^{125}I -APF-1/Ubiquitin was incubated along with crude Fraction II (Figure 3 and text) in the absence (lane 1) or presence (lanes 2-5) of ATP and in the absence (lanes 1, 2) or presence (lanes 3-5) of increasing concentrations of unlabeled lysozyme. Reaction mixtures resolved in lanes 6 and 7 were incubated in the absence (lane 6) or presence (lane 7) of ATP, and included unlabeled APF-1/ubiquitin and ^{125}I -labeled lysozyme. C1-C6 denote specific APF-1/ubiquitin-lysozyme adducts in which the number of APF-1/ubiquitin moieties bound to the lysozyme moiety of the adduct is increasing, probably from 1 to 6. Reaction mixtures were resolved via sodium dodecyl sulfate-polyacrylamide gel electrophoresis (SDS-PAGE) and visualized following exposure to an X-ray film (autoradiography) (with permission from Proceedings of the National Academy of the USA. Published originally in Ref. 40).

tified, protease that cannot recognize the unmodified substrate; following degradation, reusable APF-1 was released (40). Amino-acid analysis of APF-1, along with its known molecular mass and other general characteristics raised the suspicion that APF-1 was ubiquitin (41), a known protein of previously unknown function. Indeed, Wilkinson and colleagues confirmed unequivocally that APF-1 was indeed ubiquitin (42). Ubiquitin had been first described as a small, heat-stable and highly evolutionarily conserved protein of 76 residues. It was first purified during the isolation of thymopoietin (43) and was subsequently found to be ubiquitously expressed in all kingdoms of living cells, including prokaryotes (44). Interestingly, it was initially found to have lymphocyte-differentiating properties, a characteristic that was attributed to the stimulation of adenylate cyclase (44, 45). Accordingly, it was named UBIP for UBiquitous Immunopoietic Polypeptide (44). However, later studies showed that ubiquitin was not involved in the immune response (46), and that it was a contaminating endotoxin in the preparation that generated the adenylate cyclase and the T-cell differentiating activities. Furthermore, the sequence of several eubacteria and archaebacteria genomes as well as biochemical analyses in these organisms (unpublished) showed that ubiquitin was restricted only to eukaryotes. The finding of ubiquitin in bacteria (44) was probably due to contamination of the bacterial extract with yeast ubiquitin derived from the yeast extract in which the bacteria were grown. While in retrospect the name ubiquitin is a misnomer, as it is restricted to eukaryotes and is not ubiquitous as was previously thought, it has remained the name of the protein. The reason is probably because it was the name that was first assigned to the protein, and scientists and nomenclature committees tend, in general, to respect this tradition. Accordingly, and in order to avoid confusion, I suggest that the names of other novel enzymes and components of the ubiquitin system, but also of other systems as well, should remain as were first coined by their discoverers.

An important development in the ubiquitin research field was the discovery that a single ubiquitin moiety can be covalently conjugated to histones, particularly to histones H2A and H2B. While the function of these adducts has remained elusive until recently, their structure was unraveled in the mid-1970s. The structure of the ubiquitin conjugate of H2A (uH2A; was also designated protein A24) was deciphered by Goldknopf and Busch (47, 48) and by Hunt and Dayhoff (49) who

found that the two proteins are linked through a fork-like, branched isopeptide bond between the carboxy-terminal glycine of ubiquitin (Gly⁷⁶) and the ϵ -NH₂ group of an internal lysine (Lys¹¹⁹) of the histone molecule. The isopeptide bond found in the histone-ubiquitin adduct was suggested to be identical to the bond that was found between ubiquitin and the target proteolytic substrate (50), and between the ubiquitin moieties in the polyubiquitin chain (51,52) that was synthesized on the substrate and that functions as a proteolysis recognition signal for the downstream 26S proteasome. In this particular polyubiquitin chain the linkage is between Gly⁷⁶ of one ubiquitin moiety and internal Lys⁴⁸ of the previously conjugated moiety. Only Lys⁴⁸-based ubiquitin chains are recognized by the 26S proteasome and serve as proteolytic signals. In recent years it has been shown that the first ubiquitin moiety can also be attached in a linear mode to the N-terminal residue of the proteolytic target substrate (53). However, the subsequent ubiquitin moieties are generating Lys⁴⁸-based polyubiquitin chain on the first linearly fused moiety. N-terminal ubiquitination is clearly required for targeting naturally occurring lysine-less proteins for degradation. Yet, several lysine-containing proteins have also been described that traverse this pathway, the muscle-specific transcription factor MyoD for example. In these proteins the internal lysine residues are probably not accessible to the cognate ligases. Other types of polyubiquitin chains have also been described that are not involved in targeting the conjugated substrates for proteolysis. Thus, a Lys⁶³-based polyubiquitin chain has been described that is probably necessary to activate transcription factors (reviewed recently in Ref. 54). Interestingly, the role of monoubiquitination of histones has also been identified recently, and this modification is also involved in regulation of transcription, probably via modulation of the structure of the nucleosomes (for recent reviews, see, for example, Refs. 55, 56).

The identification of APF-1 as ubiquitin, and the discovery that a high-energy isopeptide bond, similar to the one that links ubiquitin to histone H2A, links it also to the target proteolytic substrate, resolved at that time the enigma of the energy requirement for intracellular proteolysis (see however below) and paved the road to the untangling of the complex mechanism of isopeptide bond formation. This process turned out to be similar to that of peptide bond formation that is catalyzed by tRNA synthetase following amino acid activation during protein synthesis or during the non-ribosomal synthesis of short peptides

(57). Using the unraveled mechanism of ubiquitin activation and immobilized ubiquitin as a »covalent« affinity bait, the three enzymes that are involved in the cascade reaction of ubiquitin conjugation were purified by Ciechanover, Hershko, and their colleagues. These enzymes are: (i) E1, the ubiquitin-activating enzyme, (ii) E2, the ubiquitin-carrier protein, and (iii) E3, the ubiquitin-protein ligase (58, 59). The discovery of an E3 which was a specific substrate-binding component, indicated a possible solution to the problem of the varying stabilities of different proteins – they might be specifically recognized and targeted by different ligases.

In a short period, the ubiquitin tagging hypothesis received substantial support. For example, Chin and colleagues injected into HeLa cells labeled ubiquitin and hemoglobin and denatured the injected hemoglobin by oxidizing it with phenylhydrazine. They found that ubiquitin conjugation to globin was markedly enhanced by denaturation of hemoglobin and the concentration of globin-ubiquitin conjugates was proportional to the rate of hemoglobin degradation (60). Hershko and colleagues observed a similar correlation for abnormal, amino acid analogue-containing short-lived proteins (61). A previously isolated cell cycle arrest mutant that loses the ubiquitin-histone H2A adduct at the permissive temperature (62), was found by Finley, Ciechanover and Varshavsky to harbor a thermolabile E1 (63). Following heat inactivation, the cells fail to degrade normal short-lived proteins (64). Although the cells did not provide direct evidence for substrate ubiquitination as a destruction signal, they still provided the strongest direct linkage between ubiquitin conjugation and degradation.

At this point, the only missing link was the identification of the downstream protease that would specifically recognize ubiquitinated substrates. Tanaka and colleagues identified a second ATP-requiring step in the reticulocyte proteolytic system, which occurred after ubiquitin conjugation (65), and Hershko and colleagues demonstrated that the energy was required for conjugate degradation (66). An important advance in the field was a discovery by Hough and colleagues, who partially purified and characterized a high-molecular mass alkaline protease that degraded ubiquitin adducts of lysozyme but not untagged lysozyme, in an ATP-dependent mode (67). This protease which was later called the 26S proteasome (see below), provided all the necessary criteria for being the specific proteolytic arm of the ubiquitin system. This finding was confirmed, and the protease was further

characterized by Waxman and colleagues who found that it was an unusually large, ~1.5 MDa enzyme, unlike any other known protease (68). A further advance in the field was the discovery (69) that a smaller neutral multi-subunit 20S protease complex that was discovered together with the larger 26S complex, was similar to a »multicatalytic proteinase complex« (MCP) that had been described earlier in bovine pituitary gland by Wilk and Orłowski (70). This 20S protease was ATP-independent and has different catalytic activities, cleaving on the carboxy-terminal side of hydrophobic, basic and acidic residues. Hough and colleagues raised the possibility – although they did not show it experimentally – that this 20S protease could be a part of the larger 26S protease that degrades the ubiquitin adducts (69). Later studies showed that indeed, the 20S complex is the core catalytic particle of the larger 26S complex (71, 72). However, a strong evidence that the active »mushroom«-shaped 26S protease was generated through the assembly of two distinct sub-complexes – the catalytic 20S cylinder-like MCP and an additional 19S ball-shaped sub-complex (that was predicted to have a regulatory role) – was provided only in the early 1990s by Hoffman and colleagues (73) who mixed the two purified particles and generated the active 26S enzyme.

The proteasome is a large, 26S, multicatalytic protease that degrades polyubiquitinated proteins to small peptides. It is composed of two sub-complexes: a 20S core particle (CP) that carries the catalytic activity, and a regulatory 19S regulatory particle (RP). The 20S CP is a barrel-shaped structure composed of four stacked rings, two identical outer α rings and two identical inner β rings. The eukaryotic α and β rings are composed each of seven distinct subunits, giving the 20S complex the general structure of $\alpha_{1-7}\beta_{1-7}\beta_{1-7}\alpha_{1-7}$. The catalytic sites are localized to some of the β subunits. Each extremity of the 20S barrel can be capped by a 19S RP each composed of 17 distinct subunits, 9 in a »base« sub-complex, and 8 in a »lid« sub-complex. One important function of the 19S RP is to recognize ubiquitinated proteins and other potential substrates of the proteasome. Several ubiquitin-binding subunits of the 19S RP have been identified, although their biological roles and mode of action have not been discerned. A second function of the 19S RP is to open an orifice in the α ring that will allow entry of the substrate into the proteolytic chamber. Also, since a folded protein would not be able to fit through the narrow proteasomal channel, it was assumed that the 19S particle unfolds substrates and inserts them

into the 20S CP. Both the channel opening function and the unfolding of the substrate require metabolic energy, and indeed, the 19S RP »base« contains six different ATPase subunits. Following degradation of the substrate, short peptides derived from the substrate are released, as well as reusable ubiquitin (for a scheme describing the ubiquitin system, see Figure 5; for the structure of the 26S proteasome, see Figure 6).

Concluding remarks

The evolution of proteolysis as a centrally important regulatory mechanism has served as a remarkable example for the evolution of a novel biological concept and the accompanying battles to change paradigms. The five decades' journey between the early 1940s and early 1990s began with fierce discussions on whether cellular proteins are static as has been thought for a long time, or are turning over. The discovery of the dynamic state of proteins was followed by the discovery of the lysosome that was believed – between the mid-1950s and mid-1970s – to be the organelle within which intracellular proteins are destroyed. Independent lines of experimental evidence gradually eroded the lysosomal hypothesis and resulted in a new idea that the regulated degradation of intracellular proteins under basal metabolic conditions was mediated by a non-lysosomal machinery. This resulted in the discovery of the ubiquitin system in the late 1970s and early 1980s. Interestingly, modifications of different target substrates by ubiquitin and ubiquitin-like proteins are now known to be involved in all aspects of lysosomal degradation, such as in the generation of the autophagic vacuoles, and in the routing of cargo-carrying vesicles to the lysosome (see below). Modifications by ubiquitin and ubiquitin-like proteins are now viewed, much like phosphorylation, as a mechanism to generate recognition elements *in trans* on target proteins to which downstream effectors bind. In one case, generation of Lys⁴⁸-based polyubiquitin chains, the binding effector is the 26S proteasome that degrades the ubiquitin-tagged protein. In many other cases, different modifications serve numerous proteolytic (lysosomal) and non-proteolytic functions, such as routing of proteins to their subcellular destinations. We were fortunate at the beginning of our studies to have in mind a clear distinction between lysosomal and non-lysosomal proteolytic systems not knowing what we know nowadays that the two processes are

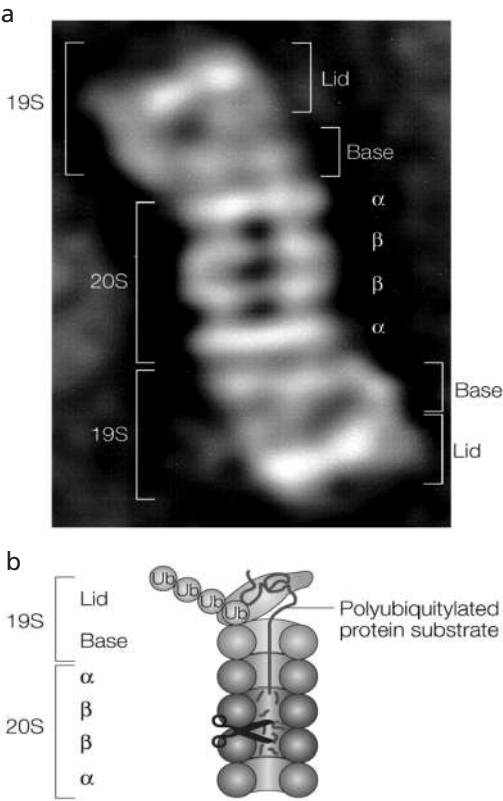
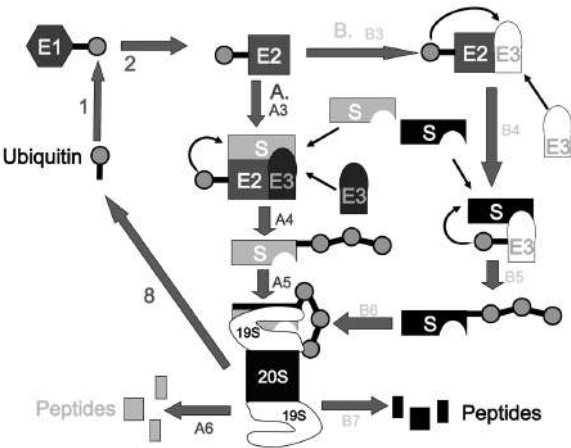


Figure 5: The ubiquitin-proteasome proteolytic system. Ubiquitin is activated by the ubiquitin-activating enzyme, E1 (1) followed by its transfer to a ubiquitin-carrier protein (ubiquitin-conjugating enzyme, UBC), E2 (2). E2 transfers the activated ubiquitin moieties to the protein substrate that is bound specifically to a unique ubiquitin ligase E3 (A and B). In the case of RING finger ligases, the transfer is direct (A3). Successive conjugation of ubiquitin moieties to one another generates a polyubiquitin chain (A4) that serves as the binding (A5) signal for the downstream 26S proteasome that degrades the target substrates to peptides (A6). In the case of HECT domain ligases, ubiquitin generates an additional thiol-ester intermediate on the ligase (B3), and only then is transferred to the substrate (B4). Successive conjugation of ubiquitin moieties to one another generates a polyubiquitin chain (B5) that binds to the 26S proteasome (B6) followed by degradation of the substrate to peptides (B7). Free and reusable ubiquitin is released by de-ubiquitinating enzymes (DUBs)(8).

Figure 6: The Proteasome. The proteasome is a large, 26S, multicatalytic protease that degrades polyubiquitinated proteins to small peptides. It is composed of two sub-complexes: a 20S core particle (CP) that carries the catalytic activity, and a regulatory 19S regulatory particle (RP). The 20S CP is a barrel-shaped structure composed of four stacked rings, two identical outer α rings and two identical inner β rings. The eukaryotic α and β rings are composed each of seven distinct subunits, giving the 20S complex the general structure of $\alpha_{1-7}\beta_{1-7}\beta_{1-7}\alpha_{1-7}$. The catalytic sites are localized to some of the β subunits. Each extremity of the 20S barrel can be capped by a 19S RP each composed of 17 distinct subunits, 9 in a »base« sub-complex, and 8 in a »lid« sub-complex. One important function of the 19S RP is to recognize ubiquitinated proteins and other potential substrates of the proteasome. Several ubiquitin-binding subunits of the 19S RP have been identified, however, their biological roles and mode of action have not been discerned. A second function of the 19S RP is to open an orifice in the α ring that will allow entry of the substrate into the proteolytic chamber. Also, since a folded protein would not be able to fit through the narrow proteasomal channel, it is assumed that the 19S particle unfolds substrates and inserts them into the 20S CP. Both the channel opening function and the unfolding of the substrate require metabolic energy, and indeed, the 19S RP »base« contains six different ATPase subunits. Following degradation of the substrate, short peptides derived from the substrate are released, as well as reusable ubiquitin (with permission from Nature Publishing Group. Published originally in Ref. 83).

- a. Electron microscopy image of the 26S proteasome from the yeast *S. cerevisiae*.
- b. Schematic representation of the structure and function of the 26S proteasome.

linked to one another and are mediated via similar modifications. Had we known that, our route would have been much more complicated.

With the identification of the reactions and enzymes that are involved in the ubiquitin-proteasome cascade, a new era in the protein degradation field began at the late 1980s and early 1990s. Studies that showed that the system was involved in targeting of key regulatory proteins – such as light-regulated proteins in plants, transcriptional factors, cell cycle regulators and tumor suppressors and promoters – started to emerge (see for example Refs. 74–78). They were followed by numerous studies on the underlying mechanisms involved in the degradation of specific proteins, each with its own unique mode of recognition and regulation. The unraveling of the human genome revealed the existence of hundreds of distinct E3s, attesting to the complexity and the high specificity and selectivity of the system. Two important advances in the field were the discovery of the non-proteolytic functions of ubiquitin such as activation of transcription and routing of proteins to the vacuole, and the discovery of modification by ubiquitin-like proteins (UBLs), that are also involved in numerous non-proteolytic functions such as directing proteins to their sub-cellular destination, protecting proteins from ubiquitination, or controlling entire processes such as autophagy (see for example Ref. 79; for the different roles of modifications by ubiquitin and UBLs, see Figure 7). All these studies have led to the emerging realization that this novel mode of covalent conjugation plays a key role in regulating a broad array of cellular process – among them cell cycle and division, growth and differentiation, activation and silencing of transcription, apoptosis, the immune and inflammatory response, signal transduction, receptor mediated endocytosis, various metabolic pathways, and the cell quality control – through proteolytic and non-proteolytic mechanisms. The discovery that ubiquitin modification plays a role in routing proteins to the lysosome/vacuole and that modification by specific and unique ubiquitin-like proteins and modification system controls autophagy closed an exciting historical cycle, since it demonstrated that the two apparently distinct systems communicate with one another. With the many processes and substrates targeted by the ubiquitin pathway, it has not been surprising to find that aberrations in the system underlie, directly or indirectly, the pathogenesis of many diseases. While inactivation of a major enzyme such as E1 was obviously lethal, mutations in enzymes or in recognition motifs in substrates that do not affect vital

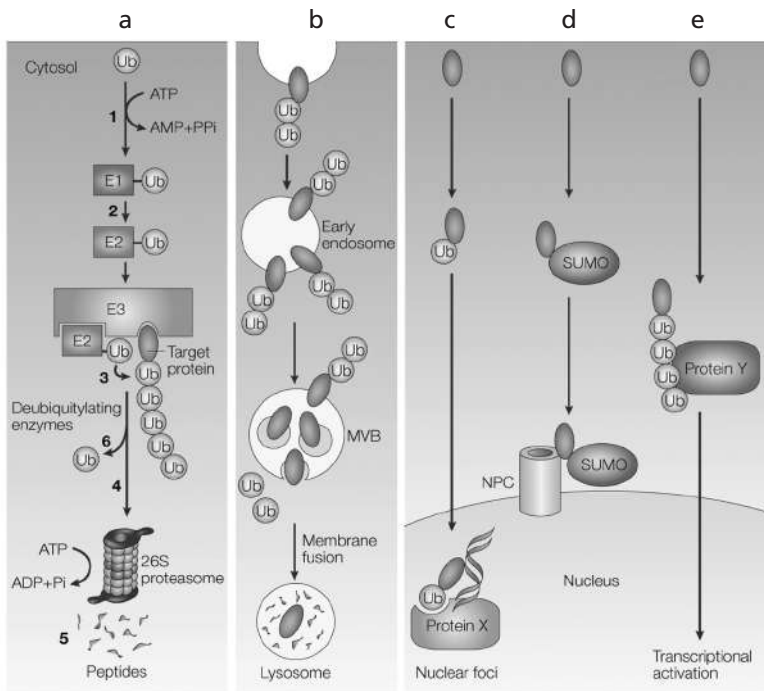


Figure 7: Some of the different functions of modification by ubiquitin and ubiquitin-like proteins.

- Proteasomal-dependent degradation of cellular proteins (see Figure 4).
- Mono- or oligoubiquitination targets membrane proteins to degradation in the lysosome/vacuole.
- Monoubiquitination, or
- a single modification by a ubiquitin-like (UBL) protein, SUMO for example, can target proteins to different subcellular destinations such as nuclear foci or the nuclear pore complex (NPC). Modification by UBLs can serve other, non-proteolytic, functions, such as protecting proteins from ubiquitination or activation of E3 complexes.
- Generation of a Lys⁶³-based polyubiquitin chain can activate transcriptional regulators, directly or indirectly [via recruitment of other proteins (Protein Y is shown), or activation of upstream components such as kinases]. Ub denotes ubiquitin.

(With permission from Nature Publishing Group. Published originally in Ref. 83)

pathways or that affect the involved process only partially, may result in a broad array of phenotypes. Likewise, acquired changes in the activity of the system can also evolve into certain pathologies. The path-

ological states associated with the ubiquitin system can be classified into two groups: (a) those that result from loss of function – mutation in a ubiquitin system enzyme or in the recognition motif in the target substrate that result in stabilization of certain proteins, and (b) those that result from gain of function – abnormal or accelerated degradation of the protein target (for aberrations in the ubiquitin system that result in disease states, see Figure 8). Studies that employ targeted inactivation of genes coding for specific ubiquitin system enzymes and substrates in animals can provide a more systematic view into the broad spectrum of pathologies that may result from aberrations in ubiquitin-mediated proteolysis. Better understanding of the processes and identification of the components involved in the degradation of key regulatory proteins will lead to the development of mechanism-based drugs that will target specifically only the involved proteins. While the first drug, a specific proteasome inhibitor is already on the market (80), it appears that one important hallmark of the new era we are entering now will be the discovery of novel drugs based on targeting of specific processes such as inhibiting aberrant Mdm2- or E6-AP-mediated accelerated targeting of the tumor suppressor p53 which will lead to regain of its lost function.

Many reviews have been published on different aspects of the ubiquitin system. The purpose of this article was to bring to the reader several milestones along the historical pathway along which the ubiquitin system has been evolved. For additional reading on the ubiquitin system, the reader is referred to numerous review articles written on the subject (for some older reviews, see for example Refs. 81, and 82). Some parts of this review, including several Figures, are based on another published review article (Ref. 83).

Acknowledgement

Research in the laboratory of Aaron Ciechanover is supported by grants from the Dr. Miriam and Sheldon G. Adelson Medical Research Foundation (AMRF), the Israel Science Foundation (ISF), the German-Israeli Foundation (GIF) for Scientific Research and Development, an Israel Cancer Research Fund (ICRF) USA Professorship, the German-Israeli Project Cooperation Program (DIP), and the Rubicon European Union (EU) Network of Excellence.

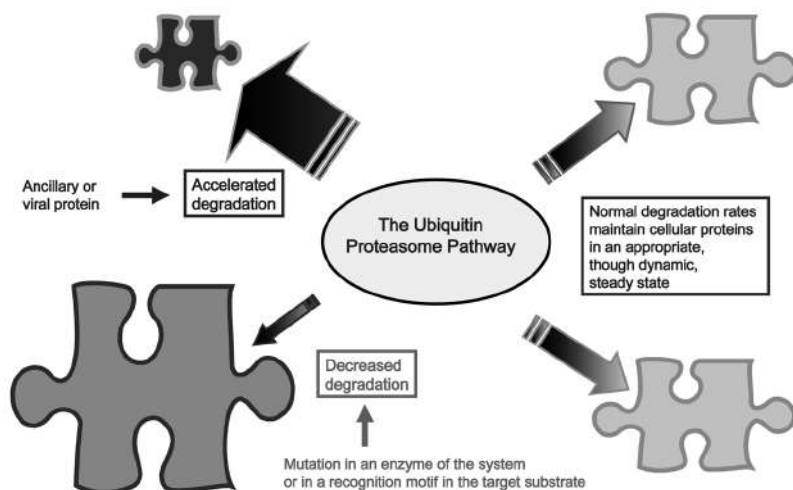


Figure 8: Aberrations in the ubiquitin-proteasome system and pathogenesis of human diseases. Normal degradation of cellular proteins maintains them in a steady state level, though this level may change under various pathophysiological conditions (upper and lower right side). When degradation is accelerated due to an increase in the level of an E3 (Skp2 in the case of p27, for example), or overexpression of an ancillary protein that generates a complex with the protein substrate and targets it for degradation (the Human Papillomavirus E6 oncoprotein that associates with p53 and targets it for degradation by the E6-AP ligase, or the cytomegalovirus-encoded ER proteins US2 and US11 that target MHC class I molecules for ERAD), the steady state level of the protein decreases (upper left side). A mutation in a ubiquitin ligase (such as occurs in Adenomatous Polyposis Coli – APC, or in E6-AP (Angelmans' Syndrome)) or in the substrate's recognition motif (such as occurs in β -catenin or in ENaC) will result in decreased degradation and accumulation of the target substrate.

References

1. Clarke, H.T. (1958). Impressions of an organic chemist in biochemistry. *Annu. Rev. Biochem.* 27, 1-14.
2. Kennedy, E.P. (2001). Hitler's gift and the era of biosynthesis. *J. Biol. Chem.* 276, 42619-42631.
3. Simoni, R.D., Hill, R.L., and Vaughan, M. (2002). The use of isotope tracers to study intermediary metabolism: Rudolf Schoenheimer. *J. Biol. Chem.* 277 (issue 43), e1-e3 (available on-line at: <http://www.jbc.org>).

4. Schoenheimer, R., Ratner, S., and Rittenberg, D. (1939). Studies in protein metabolism: VII. The metabolism of tyrosine. *J. Biol. Chem.* 127, 333-344.
5. Ratner, S., Rittenberg, D., Keston, A.S., and Schoenheimer, R. (1940). Studies in protein metabolism: XIV. The chemical interaction of dietary glycine and body proteins in rats. *J. Biol. Chem.* 134, 665-676.
6. Schoenheimer, R. *The Dynamic State of Body Constituents* (1942). Harvard University Press, Cambridge, Massachusetts, USA.
7. Hogness, D.S., Cohn, M., and Monod, J. (1955). Studies on the induced synthesis of β -galactosidase in *Escherichia coli*: The kinetics and mechanism of sulfur incorporation. *Biochim. Biophys. Acta* 16, 99-116.
8. de Duve, C., Gianetto, R., Appelmans, F., and Wattiaux, R. (1953). Enzymic content of the mitochondria fraction. *Nature (London)* 172, 1143-1144.
9. Gianetto, R., and de Duve, C. Tissue fractionation studies 4. (1955). Comparative study of the binding of acid phosphatase, β -glucuronidase and cathepsin by rat liver particles. *Biochem. J.* 59, 433-438.
10. Simpson, M.V. The release of labeled amino acids from proteins in liver slices. (1953). *J. Biol. Chem.* 201, 143-154.
11. Mortimore, G.E., and Poso, A.R. (1987). Intracellular protein catabolism and its control during nutrient deprivation and supply. *Annu. Rev. Nutr.* 7, 539-564.
12. Ashford, T.P., and Porter, K.R. (1962). Cytoplasmic components in hepatic cell lysosomes. *J. Cell Biol.* 12, 198-202.
13. Schimke, R.T., and Doyle, D. (1970). Control of enzyme levels in animal tissues. *Annual Rev. Biochem.* 39, 929-976.
14. Goldberg, A.L., and St. John, A.C. (1976). Intracellular protein degradation in mammalian and bacterial cells: Part 2. *Annu. Rev. Biochem.* 45, 747-803.
15. Segal, H.L., Winkler, J.R., and Miyagi, M.P. (1974). Relationship between degradation rates of proteins *in vivo* and their susceptibility to lysosomal proteases. *J. Biol. Chem.* 249, 6364-6365.
16. Haider, M., and Segal, H.L. (1972). Some characteristics of the alanine-aminotransferase and arginase-inactivating system of lysosomes. *Arch. Biochem. Biophys.* 148, 228-237.
17. Dean, R.T. Lysosomes and protein degradation. (1977). *Acta Biol. Med. Ger.* 36, 1815-1820.
18. Müller, M., Müller, H., and Holzer, H. (1981). Immunochemical studies on catabolite inactivation of phosphoenolpyruvate carboxykinase in *Saccharomyces cerevisiae*. *J. Biol. Chem.* 256, 723-727.
19. Holzer, H. (1989). Proteolytic catabolite inactivation in *Saccharomyces cerevisiae*. *Revis. Biol. Celular* 21, 305-319.
20. Majeski, A.E., and Dice, J.F. (2004). Mechanisms of chaperone-mediated autophagy. *Intl. J. Biochem. Cell Biol.* 36, 2435-2444.
21. Cuervo, A.M., and Dice, J.F. (1998). Lysosomes, a meeting point of proteins, chaperones, and proteases. *J. Mol. Med.* 76, 6-12.
22. Hayashi, M., Hiroi, Y., and Natori, Y. (1973). Effect of ATP on protein degradation in rat liver lysosomes. *Nature New Biol.* 242, 163-166.

23. Schneider, D.L. (1981). ATP-dependent acidification of intact and disrupted lysosomes: Evidence for an ATP-driven proton pump. *J. Biol. Chem.* 256, 3858-3864.
24. de Duve, C., and Wattiaux, R. Functions of lysosomes. (1966). *Annu. Rev. Physiol.* 28, 435-492.
25. Rabinovitz, M., and Fisher, J.M. (1964). Characteristics of the inhibition of hemoglobin synthesis in rabbit reticulocytes by threo- α -amino- β -chlorobutyric acid. *Biochim. Biophys. Acta.* 91, 313-322.
26. Carrell, R.W., and Lehmann, H. (1969). The unstable haemoglobin haemolytic anaemias. *Semin. Hematol.* 6, 116-132.
27. Huehns, E.R., and Bellingham, A.J. (1969). Diseases of function and stability of haemoglobin. *Br. J. Haematol.* 17, 1-10.
28. Etlinger, J.D., and Goldberg, A.L. (1977). A soluble ATP-dependent proteolytic system responsible for the degradation of abnormal proteins in reticulocytes. *Proc. Natl. Acad. Sci. USA* 74, 54-58.
29. Hershko, A., Heller, H., Ganoth, D., and Ciechanover, A. (1978). Mode of degradation of abnormal globin chains in rabbit reticulocytes. In: *Protein Turnover and Lysosome Function* (H.L. Segal & D.J. Doyle, eds.). Academic Press, New York. pp. 149-169.
30. Knowles, S.E., and Ballard, F.J. (1976). Selective control of the degradation of normal and aberrant proteins in Reuber H35 hepatoma cells. *Biochem J.* 156, 609-617.
31. Neff, N.T., DeMartino, G.N., and Goldberg, A.L. (1979). The effect of protease inhibitors and decreased temperature on the degradation of different classes of proteins in cultured hepatocytes. *J. Cell Physiol.* 101, 439-457.
32. Poole, B., Ohkuma, S., and Warburton, M.J. (1977). The accumulation of weakly basic substances in lysosomes and the inhibition of intracellular protein degradation. *Acta Biol. Med. Germ.* 36, 1777-1788.
33. Poole, B., Ohkuma, S. & Warburton, M.J. (1978). Some aspects of the intracellular breakdown of exogenous and endogenous proteins. In: *Protein Turnover and Lysosome Function* (H.L. Segal & D.J. Doyle, eds.). Academic Press, New York. pp. 43-58.
34. Mandelstam, J. (1958). Turnover of protein in growing and non-growing populations of *Escherichia coli*. *Biochem. J.* 69, 110-119.
35. Steinberg, D., and Vaughan, M. (1956). Observations on intracellular protein catabolism studied *in vitro*. *Arch. Biochem. Biophys.* 65, 93-105.
36. Hershko, A., and Tomkins, G.M. (1971). Studies on the degradation of tyrosine aminotransferase in hepatoma cells in culture: Influence of the composition of the medium and adenosine triphosphate dependence. *J. Biol. Chem.* 246, 710-714.
37. Goldberg, A.L., Kowit, J.D., and Etlinger, J.D. (1976). Studies on the selectivity and mechanisms of intracellular protein degradation. In: *Proteolysis and physiological regulation* (D.W. Ribbons & K. Brew, eds.). Academic Press, New York. pp. 313-337.

38. Ciechanover A., Hod, Y., and Hershko, A. (1978). A heat-stable polypeptide component of an ATP-dependent proteolytic system from reticulocytes. *Biochem. Biophys. Res. Commun.* 81, 1100-1105.
39. Ciechanover, A., Heller, H., Elias, S., Haas, A.L., and Hershko, A. (1980). ATP-dependent conjugation of reticulocyte proteins with the polypeptide required for protein degradation. *Proc. Natl. Acad. Sci. USA.* 77, 1365-1368.
40. Hershko, A., Ciechanover, A., Heller, H., Haas, A.L., and Rose, I.A. (1980). Proposed role of ATP in protein breakdown: Conjugation of proteins with multiple chains of the polypeptide of ATP-dependent proteolysis. *Proc. Natl. Acad. Sci. USA* 77, 1783-1786.
41. Ciechanover, A., Elias, S., Heller, H., Ferber, S. & Hershko, A. (1980). Characterization of the heat-stable polypeptide of the ATP-dependent proteolytic system from reticulocytes. *J. Biol. Chem.* 255, 7525-7528.
42. Wilkinson, K.D., Urban, M.K., and Haas, A.L. (1980). Ubiquitin is the ATP-dependent proteolysis factor I of rabbit reticulocytes. *J. Biol. Chem.* 255, 7529-7532.
43. Goldstein, G. (1974). Isolation of bovine thymosin, a polypeptide hormone of the thymus. *Nature (London)* 247, 11-14.
44. Goldstein, G., Scheid, M., Hammerling, U., Schlesinger, D.H., Niall, H.D., and Boyse, E.A. (1975). Isolation of a polypeptide that has lymphocyte-differentiating properties and is probably represented universally in living cells. *Proc. Natl. Acad. Sci. USA* 72, 11-15.
45. Schlessinger, D.H., Goldstein, G., and Niall, H.D. (1975). The complete amino acid sequence of ubiquitin, an adenylate cyclase stimulating polypeptide probably universal in living cells. *Biochemistry* 14, 2214-2218.
46. Low, T.L.K., and Goldstein, A.L. (1979). The chemistry and biology of thymosin: amino acid analysis of thymosin α 1 and polypeptide β 1. *J. Biol. Chem.* 254, 987-995.
47. Goldknopf, I.L., and Busch, H. (1975). Remarkable similarities of peptide fingerprints of histone 2A and non-histone chromosomal protein A24. *Biochem. Biophys. Res. Commun.* 65, 951-955.
48. Goldknopf, I.L., and Busch, H. (1977). Isopeptide linkage between non-histone and histone 2A polypeptides of chromosome conjugate-protein A24. *Proc. Natl. Acad. Sci. USA* 74, 864-868.
49. Hunt, L.T., and Dayhoff, M.O. (1977). Amino-terminal sequence identity of ubiquitin and the non-histone component of nuclear protein A24. *Biochim. Biophys. Res. Commun.* 74, 650-655.
50. Hershko, A., Ciechanover, A., and Rose, I.A. (1981). Identification of the active amino acid residue of the polypeptide of ATP-dependent protein breakdown. *J. Biol. Chem.* 256, 1525-1528.
51. Hershko, A., and Heller, H. (1985). Occurrence of a polyubiquitin structure in ubiquitin-protein conjugates. *Biochem. Biophys. Res. Commun.* 128, 1079-1086.
52. Chau, V., Tobias, J. W., Bachmair, A., Mariott, D., Ecker, D., Gonda, D. K., and Varshavsky, A. (1989). A multiubiquitin chain is confined

- to specific lysine in a targeted short-lived protein. *Science* 243, 1576-1583.
53. Ciechanover, A., and Ben-Saadon R. (2004). N-terminal ubiquitination: More Protein substrates join in. *Trends Cell Biol.* 14, 103-106.
 54. Muratani, M., and Tansey, W.P. (2003). How the ubiquitin-proteasome system controls transcription. *Nat. Rev. Mol. Cell Biol.* 4, 192-201.
 55. Zhang, Y. (2003). Transcriptional regulation by histone ubiquitination and deubiquitination. *Genes & Dev.* 17, 2733-2740.
 56. Osley, M.A. (2004). H2B ubiquitylation: the end is in sight. *Biochim. Biophys. Acta.* 1677, 74-78.
 57. Lipman, F. (1971). Attempts to map a process evolution of peptide biosynthesis. *Science* 173, 875-884.
 58. Ciechanover, A., Elias, S., Heller, H. & Hershko, A. (1982). »Covalent affinity« purification of ubiquitin-activating enzyme. *J. Biol. Chem.* 257, 2537-2542.
 59. Hershko, A., Heller, H., Elias, S., and Ciechanover, A. (1983). Components of ubiquitin-protein ligase system: Resolution, affinity purification and role in protein breakdown. *J. Biol. Chem.* 258, 8206-8214.
 60. Chin, D.T., Kuehl, L., and Rechsteiner, M. (1982). Conjugation of ubiquitin to denatured hemoglobin is proportional to the rate of hemoglobin degradation in HeLa cells. *Proc. Natl. Acad. Sci. USA* 79, 5857-5861.
 61. Hershko, A., Eytan, E., Ciechanover, A. and Haas, A.L. (1982). Immunochemical Analysis of the Turnover of Ubiquitin-protein Conjugates in Intact Cells: Relationship to the Breakdown of Abnormal Proteins. *J. Biol. Chem.* 257, 13 964-13 970.
 62. Matsumoto, Y., Yasuda, H., Marunouchi, T., and Yamada, M. (1983). Decrease in uH2A (protein A24) of a mouse temperature-sensitive mutant. *FEBS Lett.* 151, 139-142.
 63. Finley, D., Ciechanover, A., and Varshavsky, A. (1984). Thermolability of ubiquitin-activating enzyme from the mammalian cell cycle mutant ts85. *Cell* 37, 43-55.
 64. Ciechanover, A., Finley D., and Varshavsky, A. (1984). Ubiquitin dependence of selective protein degradation demonstrated in the mammalian cell cycle mutant ts85. *Cell* 37, 57-66.
 65. Tanaka, K., Waxman, L., and Goldberg, A.L. (1983). ATP serves two distinct roles in protein degradation in reticulocytes, one requiring and one independent of ATP. *J. Cell Biol.* 96, 1580-1585.
 66. Hershko, A., Leshinsky, E., Ganoth, D. & Heller, H. (1984). ATP-dependent degradation of ubiquitin-protein conjugates. *Proc. Natl. Acad. Sci. USA* 81, 1619- 1623.
 67. Hough, R., Pratt, G. & Rechsteiner, M. (1986). Ubiquitin-lysozyme conjugates. Identification and characterization of an ATP-dependent protease from rabbit reticulocyte lysates. *J. Biol. Chem.* 261, 2400-2408.
 68. Waxman, L., Fagan, J., and Goldberg, A.L. (1987). Demonstration of two distinct high molecular weight proteases in rabbit reticulocytes, one of which degrades ubiquitin conjugates. *J. Biol. Chem.* 262, 2451-2457.

69. Hough, R., Pratt, G., and Rechsteiner M. (1987). Purification of two high molecular weight proteases from rabbit reticulocyte lysate. *J. Biol. Chem.* 262, 8303-8313.
70. Wilk, S., and Orłowski, M. (1980). Cation-sensitive neutral endopeptidase: isolation and specificity of the bovine pituitary enzyme. *J. Neurochem.* 35, 1172-1182.
71. Eytan, E., Ganoth, D., Armon, T., and Hershko, A. (1989). ATP-dependent incorporation of 20S protease into the 26S complex that degrades proteins conjugated to ubiquitin. *Proc. Natl. Acad. Sci. USA.* 86, 7751-7755.
72. Driscoll, J., and Goldberg, A.L. (1990). The proteasome (multicatalytic protease) is a component of the 1,500-kDa proteolytic complex which degrades ubiquitin-conjugated proteins. *J. Biol. Chem.* 265, 4789-4792.
73. Hoffman, L., Pratt, G., and Rechsteiner, M. (1992). Multiple forms of the 20S multicatalytic and the 26S ubiquitin/ATP-dependent proteases from rabbit reticulocyte lysate. *J. Biol. Chem.* 267, 22362-22368.
74. Shanklin, J., Jaben, M., and Vierstra, R.D. (1987). Red light-induced formation of ubiquitin-phytochrome conjugates: Identification of possible intermediates of phytochrome degradation. *Proc. Natl. Acad. Sci. USA* 84, 359-363.
75. Hochstrasser, M., and Varshavsky, A. (1990). *In vivo* degradation of a transcriptional regulator: the yeast $\alpha 2$ repressor. *Cell* 61, 697-708.
76. Scheffner, M., Werner, B.A., Huibregtse, J.M., Levine, A.J., and Howley, P.M. (1990). The E6 oncoprotein encoded by human papillomavirus types 16 and 18 promotes the degradation of p53. *Cell* 63, 1129-1136.
77. Glotzer, M., Murray, A.W., and Kirschner, M.W. (1991). Cyclin is degraded by the ubiquitin pathway. *Nature* 349, 132-138.
78. Ciechanover, A., DiGiuseppe, J.A., Bercovich, B., Orian, A., Richter, J.D., Schwartz, A.L., and Brodeur, G.M. (1991). Degradation of nuclear oncoproteins by the ubiquitin system *in vitro*. *Proc. Natl. Acad. Sci. USA* 88, 139-143.
79. Mizushima, N., Noda, T., Yoshimori, T., Tanaka, Y., Ishii, T., George, M.D., Klionsky, D.J., Ohsumi, M., and Ohsumi, Y. (1998). A protein conjugation system essential for autophagy. *Nature* 395, 395-398.
80. Adams J. (2003). Potential for proteasome inhibition in the treatment of cancer. *Drug Discov. Today.* 8, 307-315.
81. Glickman, M.H., and Ciechanover, A. (2002). The ubiquitin-proteasome pathway: Destruction for the sake of construction. *Physiological Reviews* 82, 373-428.
82. Pickart, C.M., and Cohen, R.E. (2004). Proteasomes and their kin: proteases in the machine age. *Nature Rev. Mol. Cell Biol.* 5, 177-187.
83. Ciechanover, A. (2005). From the lysosome to ubiquitin and the proteasome. *Nature Rev. Mol. Cell Biol.* 6, 79-86.

Eckart Altenmüller, Oliver Grewe,
Frederik Nagel und Reinhard Kopiez

Musik als Sprache der Gefühle

Neurobiologische und musikpsychologische Aspekte¹

Musizieren und Musikhören wird nach Umfragen von der Mehrzahl der Deutschen immer noch als die wichtigste Freizeitbeschäftigung angesehen. Immerhin etwa 4 Millionen Mitbürger musizieren regelmäßig an einem Instrument oder singen in einem Chor (Gembris 1998). Die neurobiologischen und anthropologischen Grundlagen dieser Vorliebe für Musik sind bislang wenig erforscht, aber allgemein wird angenommen, dass die Wirkung der Musik vor allem auf der Auslösung intensiver Emotionen beruht. Häufig wird daher Musik auch als Sprache der Gefühle bezeichnet.

Bevor wir uns der Frage zuwenden, welche Musik bei Menschen starke Emotionen auslösen kann, sollen die Begriffe Musik und Emotion definiert werden. Musik sind bewusst gestaltete, zeitlich strukturierte akustische Phänomene. In dieser reduktionistischen Musikdefinition fehlt allerdings noch eine wesentliche Eigenschaft von Musik, nämlich ihr kommunikativer Aspekt. Man sollte also eher sagen, Musik ist die bewusst gestaltete, zeitlich strukturierte Ordnung von nicht sprachlichen akustischen Ereignissen in sozialen Kontexten. Als Argument für diese erweiterte Definition kann angeführt werden, dass Musik in zahlreichen sozialen Kontexten stattfindet und häufig spezifische Funktionen erfüllt. Viel zitierte Beispiele sind die Wiegenlieder, die der Mutter-Kind-Bindung und wahrscheinlich auch dem Spracherwerb dienen (Trehub 2003) oder der Tanz, der eine Ausschüttung des Hypophysenhormons Oxytocin bewirkt und dadurch die Erinnerung an ein spezifisches Gruppenerlebnis fördert. Bei zahlreichen Ge-

1 Altenmüller E., Grewe O., Nagel F., Kopiez R. (2006) Musik als Sprache der Gefühle – Neurobiologische und musikpsychologische Aspekte. Jahrbuch 2005. Leopoldina (R. 3) 51, S. 459-465. Wiederabdruck mit freundlicher Genehmigung der Deutschen Akademie der Naturforscher Leopoldina.

legenheiten wird Musik als Markersignal von Gruppenidentität eingesetzt. Man denke nur an Nationalhymnen, Fußballgesänge und an die identitätsstiftende Wirkung, die bestimmte Lieder von ethnischen Minderheiten in einem Staatswesen haben (Kopiez und Brink 1998). Schließlich dient Musik als Mittel zur Verhaltenssynchronisation, beispielsweise als Marschmusik beim Militär (McNeill 1995).

Der zweite zu definierende Begriff ist der der Emotion. Wir wollen an dieser Stelle auf eine ausführliche Terminologiediskussion verzichten und uns an die Lehrbuchdefinition halten: Beim Menschen versteht man unter Emotion »ein Reaktionsmuster, das auf vier Ebenen wirksam wird: a) als subjektives Gefühl, b) als motorische Äußerung, z.B. als Ausdrucksverhalten in Mimik, Gestik, und Stimme, c) als physiologische Reaktion des autonomen Nervensystems z.B. Erhöhung der Herzschlagfrequenz und d) als kognitive Bewertung« (Zimbardo und Gerrig 1999).

Es ist nicht einfach, Emotionen objektiv zu messen, insbesondere wenn es sich um subtile ästhetische Erlebnisse beim Musikhören handelt. Allerdings gibt es ein Phänomen, das man sich hier zunutze machen und das als »Starke Emotion beim Musikhören« (SEM, *Strong Emotions of Music*) bekannt ist (Gabrielson 2001). Es handelt sich um intensive Musikerlebnisse, die zu Reaktionen des autonomen Nervensystems führen. Sie werden häufig als Gänsehaut-Erlebnisse beschrieben. Im Englischen hat sich dafür der Begriff Chill durchgesetzt, der hier im Weiteren verwendet werden soll. In den letzten Jahren untersuchten wir, welche Musik bei welchen Menschen unter welchen Bedingungen derartige Chills auslöst. Die Ergebnisse dieser Studien sollen im Folgenden zusammengefasst werden.

Wie kann man Emotionen messen?

Die erste Schwierigkeit war, die Entwicklung von Emotionen in der Zeit zu erfassen. Um dieses Problem zu lösen, wurde ein Versuchsaufbau entwickelt, mit dem die Probanden kontinuierlich während des Hörens von Musik ihre wechselnde emotionale Befindlichkeit berichten konnten (Nagel et al. 2006). Dazu bewegten sie mit der Computermaus auf dem Bildschirm einen Cursor, mit dem sie auf einem Achsenkreuz angeben konnten, wie ihnen die Musik gefiel (Valenzurteil) und ob sie die Musik als eher aufregend oder beruhigend empfanden

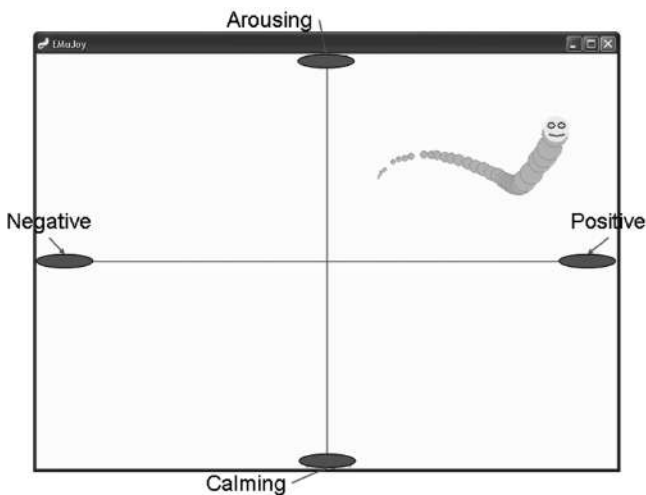


Abbildung 1: Das auf Russell (1989) beruhende zweidimensionale Emotionsmodell mit den Achsen Valenz (Gefallen) und *Arousal* (Erregung). Die Probanden konnten während des Hörens von Musik ihre momentane Stimmungslage auf diesem Achsenkreuz mit einer Computermaus angeben.

(*Arousal*-Urteil, Abb. 1). Chill-Erlebnisse wurden mit dem Drücken einer Maustaste angezeigt. Als psychophysiologische Marker des emotionalen Erlebens registrierten wir Hautleitwert, Aktivität von Emotionen anzeigenden Gesichtsmuskeln, Atemfrequenz, Hauttemperatur und Herzrate.

Während eines öffentlichen Vortrags wählten wir 38 Versuchspersonen unterschiedlichen Alters und unterschiedlicher sozialer Herkunft aus. Ihr Altersdurchschnitt war 38 Jahre (11–72 Jahre), neun Probanden waren männlich, 29 weiblich. Fünfundzwanzig Probanden spielten oder hatten früher ein Instrument gespielt, dreizehn Probanden waren musikalische Laien. Die Probanden hörten sieben zuvor von uns ausgewählte Stücke und brachten eigene Musik mit, von der sie wussten, dass sie darauf stark emotional reagieren würden. In Tabelle 1 sind die von uns ausgewählten Stücke und ihre speziellen Merkmale zusammengefasst.

Musikstück	Eigenschaften
»Tuba mirum« Requiem KV 626 Wolfgang Amadeus Mozart (Karajan, 1989)	In diesem Stück setzen die Solostimmen und der Chor nacheinander ein, ohne dass es zu Überlappungen kommt. Daher konnte der Effekt der verschiedenen Stimmlagen untersucht werden.
»Toccata BWV 540« Johann Sebastian Bach (1997)	Beispiel für ein komplexes Werk Barocker Kompositionstechnik. Unterschiedliche Themen werden entwickelt und treten wiederholt auf. Hier konnte der Effekt musikalischer Wiederholungen untersucht werden.
»Making love out of nothing« Air Supply (1997)	Pop-Musik mit einem sehr expressiven Tenor und starker emotionaler Wirkung.
»Main title« aus dem Film »Chocolat« Rachel Portman (2000)	Eher melancholische Filmmusik.
»Coma« Apocalyptica (2004)	Düster-melancholisches Musikstück der Cello-Rockband.
»Skull full of maggots« Cannibal Corpse (2002)	Sehr aggressive <i>Death metal</i> -Rockmusik. Schlagzeug und wilde Schreie dominieren dieses sehr motorische Stück.
»Bossa nova« Quincy Jones (1997)	Beschwingte und fröhliche Tanzmusik.

Tabelle 1: Musikstücke und ihre hervorstechenden Eigenschaften.

Da wir auch feststellen wollten, ob es eine typische emotionale Hörerpersönlichkeit gibt, füllten die Probanden nach dem Versuch standardisierte Persönlichkeitsskalen aus, unter anderem das Temperament- und Charakter-Inventar (TCI) von Cloninger (Cloninger et al. 1999) und die *Sensation Seeking Scale* von Litle und Zuckerman (Beauducel et al. 1999). In zwei von uns entwickelten Fragebögen wurde zusätzlich nach jedem Stück Bekanntheit, Gefallen, Angenehmheit, Erinnerungen und körperliche Reaktionen erfragt. Am Ende des Versuches wurden Fragen zu demografischen Daten, zu Erfahrung mit Musikstilen, zur musikalischen Biographie und zu Krankheiten und emotionalen Ereignissen der letzten Zeit gestellt.

Alle Chill-Musikstücke wurden nach psychoakustischen und nach musikwissenschaftlichen Kriterien analysiert. Außerdem wurden diese Eigenschaften der Musik zu den physiologischen Daten und zur Selbstauskunft in Beziehung gesetzt. Auf die Details der Methodik und der statistischen Auswertung soll hier nicht weiter eingegangen werden, wir verweisen diesbezüglich auf zwei Publikationen (Nagel et al. 2007, Grewe et al. 2007).

Die Wissenschaft von der Gänsehaut

Die Ergebnisse dieser explorativen Studie zeigten erstmalig, welche Menschen bei welcher Musik starke Emotionen erleben. Insgesamt sammelten wir 291 Chill-Stellen in 126 Musikstücken. In 117 von 126 Stücken waren Hautleitwert und Chill-Erlebnisse korreliert. Meist stieg der Hautleitwert bereits wenige Sekunden vor dem Zeitpunkt des Chills an. Ähnlich verhielt sich die Herzrate – auch sie stieg wenige Sekunden vor dem Chill-Erleben an. Die Häufigkeit der Chills war naturgemäß bei den mitgebrachten Stücken mit hohem Bekanntheitsgrad und Gefallen am größten. Überraschenderweise ist allerdings die Beziehung zwischen Chills und psychophysiologischen Antworten nicht eindeutig, d.h. es treten in seltenen Fällen Chill-Erlebnisse auf, ohne dass sich dies in psychophysiologischen Parametern widerspiegelt, umgekehrt kommt es zu starken psychophysiologischen Reaktionen, ohne dass die Probanden einen Chill angeben! Ein eindrucksvolles Beispiel ist in Abbildung 2 zu sehen.

Die psychoakustische Analyse der Chill-Stellen ergab, dass sich starke emotionale Antworten bei Zunahme der Lautheit, insbesondere

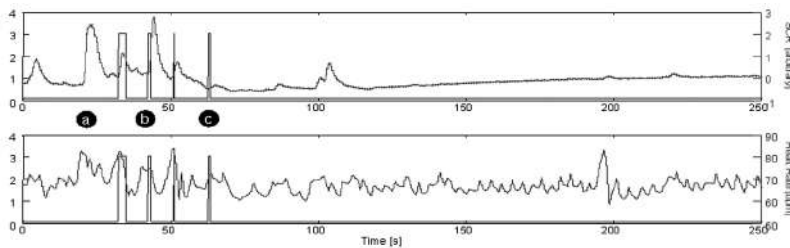


Abbildung 2: Beispiel der Ableitung von Hautleitwert (SCR, obere Kurve) und Herzrate (untere Kurve) eines Probanden. Die Chill-Antwort (Rechtecke) korreliert nur teilweise mit den beiden physiologischen Parametern. So kommt es in Situation A zu starkem Anstieg der Hautleitwerte und der Herzrate ohne Chill-Antwort. In der Situation B stimmen Chill-Antwort, Hautleitwert und Herzrate überein, in der Situation C ist die Chill-Antwort nur an eine Herzratenerhöhung gekoppelt.

im höheren Frequenzbereich um 1000 bis 3000 Hz, häuften. Abgesehen von diesen einfachen psychoakustischen Parametern entstehen starke Emotionen auch häufiger bei bestimmten musikalischen Ereignissen. Dazu gehören der Beginn eines neuen Abschnitts, der Einsatz einer Solostimme und der Einsatz eines Chores. Damit mehrten sich die Hinweise, dass einerseits emotionale Erlebnisse beim Beginn von etwas Neuem – oder anders ausgedrückt, bei einer Verletzung der Erwartung – entstehen und dass andererseits die menschliche Stimme ein besonderes Potential hat, starke Emotionen auszulösen. Eine Detailanalyse der Stücke ergibt aber auch Hinweise darauf, dass das Erkennen von Strukturen, zum Beispiel in der Toccata von Johann Sebastian Bach, starke Emotionen auslösen kann. Derartige eher intellektuelle Erlebnisse sind jedoch geübten Hörern vorbehalten, da die Fähigkeit zum Erkennen dieser komplexen Strukturen eine Vorbedingung für diese Chills zu sein scheint. Schließlich hängen starke Erlebnisse auch von bestimmten Persönlichkeitseigenschaften ab. So haben Chill-Persönlichkeiten grundsätzlich niedrigere Reizschwellen bezüglich ihres emotionalen Erlebens, sie sind tendenziell schüchterner und abhängiger von Belohnungen, sie identifizieren sich stärker mit der von ihnen bevorzugten Musik und Musik hat eine größere Wichtigkeit in ihrem Leben.

Ausblick

Wie kann man diese heterogenen Ergebnisse einordnen? In Abbildung 3 präsentieren wir ein Arbeitsmodell zu der Entstehung von Chill-Erlebnissen. Wir gehen dabei von einem mehrstufigen Prozess aus, bei dem die Qualität der Musik, etwa der Einsatz der menschlichen Stimme oder der Beginn von etwas Neuem, die Persönlichkeitseigenschaften der Hörer, z.B. Empfindsamkeit und Belohnungsabhängigkeit, und die musikalische Biographie wichtig sind. Alle drei Faktoren müssen zusammen kommen, um eine hohe Wahrscheinlichkeit eines Chill-Erlebens zu erreichen.

Es gibt auch einen direkten Weg, der Chills erzeugt, z.B. bei plötzlicher Lautstärkenänderung in dem berühmten Barrabasruf aus der Matthäuspasion von Johann Sebastian Bach. Diese Gänsehauterlebnisse hängen mit Orientierungs- und Schreckreaktionen zusammen.

Was aber sind nun die evolutionären Ursprünge dieser Chill-Antwort? Bei der Beantwortung dieser Frage muss man bedenken, dass Chill-Erlebnisse mit einer Ausschüttung von Endorphinen einhergehen, die das Belohnungssystem aktivieren und die Gedächtnisbildung verstärken (Blood und Zatorre 2001). Der biologische Hintergrund des Gänsehaut-Reflexes, der ja dem Chill-Erleben zu Grunde liegt, besteht in der Erwärmung durch Aufstellen der Körperbehaarung. Die Trennungsrufe mancher Primaten führen ebenfalls zu einer Aufstellung der Körperbehaarung und sind ein verwandtes Phänomen. Man kann also spekulieren, dass Chill-Reaktionen ursprünglich Bestandteil eines lautlichen Kommunikationssystem gewesen sind, das dazu diente, soziale Bindung – metaphorisch gesprochen »Wärme« – zu erzeugen, wichtige Veränderungen der Hörwelt anzuzeigen und durch Ausschüttung von Endorphinen die Gedächtnisbildung zu unterstützen. Später wurden möglicherweise Chill-Reaktionen spielerisch in der Musik als Mittel zur Selbstbelohnung unter feindlichen Lebensbedingungen eingesetzt.



Abbildung 3: Dreistufiges Modell der Auslösung von Chill-Antworten. Der direkte Weg führt bei stark überschwelligen, lauten oder auch aversiven Geräuschen zu einem Chill.

Literatur

- Beauducel A.B., Strobel A., Strobel A. (1999) Construct validity of sensation seeking. *Zeitschrift für differentielle und diagnostische Psychologie* 20(3), S. 155-171.
- Blood A., Zatorre R. (2001) Intensely pleasurable responses to music correlate with activity in brain regions implicated in reward and emotion. *Proceedings of the National Academy of Science* 98, S. 11 818-11 823.
- Cloninger R.C.P., Thomas R., Svrakic D., Wetzel R. (1999) Das Temperament- und Charakter-Inventar. Richter M.J.E., Richter G., Cloninger R.C.P., translator. Frankfurt (Swets & Zeitlinger BV).
- Gabrielson A. (2001) Emotions in strong experiences with music. In: *Music and emotion: Theory and research*, hrsg. v. P. Juslin, J. Sloboda, Oxford (Oxford University Press), S. 431-452.
- Gembris H. (1998) *Grundlagen musikalischer Entwicklung*. Augsburg (Wißner).
- Grewe O., Nagel F., Kopiez R., Altenmüller E. (2007) Listening to Music as a Re-Creative Process – Physiological, Psychological and Psychoacoustical Correlates of Chills and Strong Emotions. *Music Perception* 24, S. 297-315.
- Kopiez R., Brink G. (1998) *Fußball-Fangesänge. Eine Fanomenologie*. Würzburg (Königshausen & Neumann).
- McNeill W.H. (1995) *Keeping together in time. Dance and drill in human history*. Cambridge (Harvard University Press).

- Nagel F., Grewe O., Kopiez R., Altenmüller E. (2007) EMuJoy' – Software for continuous measurement of perceived emotions in music: Basic aspects of data recording and interface features. *Behavior Research Methods* 39, S. 283-290.
- Russell J. (1989) Measures of emotion. In: *Emotion: Theory research and experience* (Vol. 4), hrsg. v. R. Plutchik, H. Kellerman. New York (Academic Press), S. 81-111.
- Trehub S. (2003) Musical predispositions in infancy: an update. In: *The cognitive Neuroscience of Music*, hrsg. v. I. Peretz, R. Zatorre. Oxford (Oxford University Press), S. 57-78.
- Zimbardo P.G., Gerrig R.J. (1999) *Psychologie*. Berlin (Springer Verlag).

Birger Kollmeier

Partys, MP3-Player, Hörgeräte: Leistungen physikalischer Hörmodelle¹

gewidmet Prof. Dr. Manfred R. Schroeder (1926-2009)

Die Physik wächst an den Rändern: Während vor 20 Jahren Physiker, die sich mit neuronalen Systemen oder Sinnesorganen beschäftigten, eher zu den Außenseitern zählten, ist die Beschäftigung mit komplexen, nichtlinearen biologischen Systemen zu einem auch für den wissenschaftlichen Nachwuchs attraktiven Feld geworden, das von der Interdisziplinarität und der Anwendungsbreite physikalischer Methoden lebt. Nicht erst seit Helmholtzs frühen Arbeiten über »Die Lehre von den Tonempfindungen« (v. Helmholtz 1870) ist die Beschäftigung mit dem Hörsinn ein wichtiges Beispiel hierfür. Dieser Beitrag zeigt die Verbindung auf zwischen der medizinischen Sichtweise des Hörvorgangs (und seinen möglichen Störungen) und der physikalischen Sichtweise. Mögliche Anwendungen des quantitativen Zugangs zur Modellierung des Hörvorgangs werden exemplarisch für die Bereiche Audio-Signalkodierung im MP3-Format und digitale Hörgeräte vorgestellt.

Eine wesentliche Motivation für die Beschäftigung mit dieser Thematik ist der so genannte Cocktail-Party-Effekt (Abb. 1), der Partygästen die Konversation erleichtert und dessen Ausfall eines der schwerwiegendsten Probleme für Schwerhörige darstellt: In einer störrauschbehafteten Umgebung (z.B. einer lebhaften Party) können sich Normalhörende relativ gut auf einen Sprecher konzentrieren und die Störrausche unterdrücken. Schwerhörigen fällt dies jedoch besonders schwer, sodass sie derartige Situationen und größere Men-

¹ *Manfred-Schroeder-Lecture* beim XLAB Science Festival 2011. Der vorliegende Beitrag ist eine Überarbeitung des Artikels »Cocktail-Partys und Hörgeräte: Biophysik des Gehörs. Physikalisch inspirierte Hörmodelle weisen den Weg zu intelligenten Hörgeräten« von Birger Kollmeier. *Physik Journal 1* (2002) Nr. 4. Copyright Wiley-VCH Verlag GmbH & Co. KGaA. Reproduced with permission.



Abbildung 1: Cocktail-Party-Effekt. In der dargestellten Situation können Normalhörende die vielen Nebengeräusche unterdrücken und sich auf ein Gespräch konzentrieren. Schwerhörige haben dagegen große Probleme – eine besondere Herausforderung für die Kommunikationsakustik.

schenansammlungen meiden, was häufig zur sozialen Isolation führt. Selbst die modernsten kommerziellen, digitalen Hörgeräte können zwar in akustisch einfachen Situationen (mit nur einer stationären Störquelle, die aus einer ganz anderen Richtung als das Nutzsignal kommt) die Störgeräusche wesentlich besser unterdrücken, als es den Geräten noch vor etwa zehn Jahren möglich war, versagen aber nach wie vor in akustisch schwierigen Situationen wie etwa einer Party. Um hier Lösungsmöglichkeiten aufzuzeigen, muss das zugrundeliegende System Ohr erst einmal analysiert werden.

Aufbau und Funktionsweise des Gehörs

Das Außen- und Mittelohr (Abb. 2) leitet den Schall mit nur geringen Energieverlusten in das mit Flüssigkeit gefüllte Innenohr. Die eigentliche Umwandlung der Schallschwingungen in Nervenimpulse findet im Innenohr (Cochlea oder Hörschnecke) auf der Basilarmembran statt. Die auf der Basilarmembran angeordneten inneren und äußeren Haarzellen stellen die mechano-elektrischen Wandler dar, in denen Auslenkungen der Härchen Nervenimpulse auslösen, die zum Gehirn weitergeleitet werden.

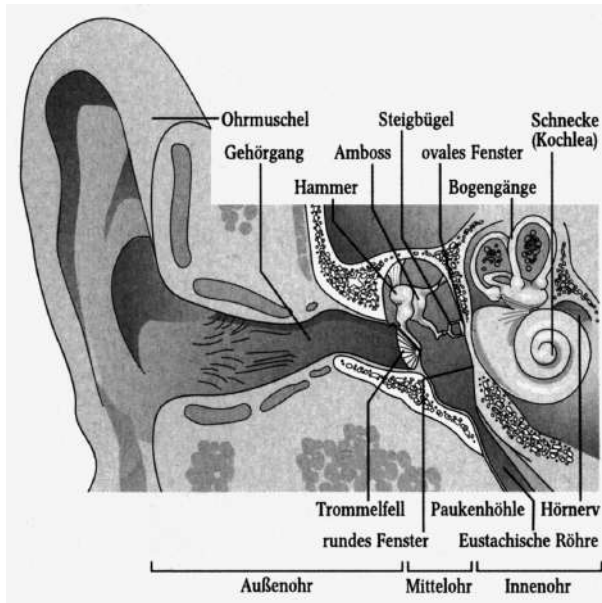


Abbildung 2: Das Ohr aus Sicht des Mediziners. Der Schall gelangt (von links nach rechts) über das Außenohr (Ohrmuschel, Gehörgang, Trommelfell) als mechanische Schwingung über die Gehörknöchelchen des Mittelohrs (Hammer, Amboss, Steigbügel) in das Innenohr (Cochlea). Dort lenken die Schwingungen je nach Frequenz verschiedene Orte der Basilarmembran maximal aus. Durch die Haarzellen auf der Basilarmembran werden ihre mechanischen Schwingungen in elektrische Spannungsimpulse im nachgeschalteten Hörnerv umgewandelt, die im zentralen auditorischen System im Gehirn (hier nicht dargestellt) ausgewertet werden (Bildnachweis: Kießling et al. 2008).

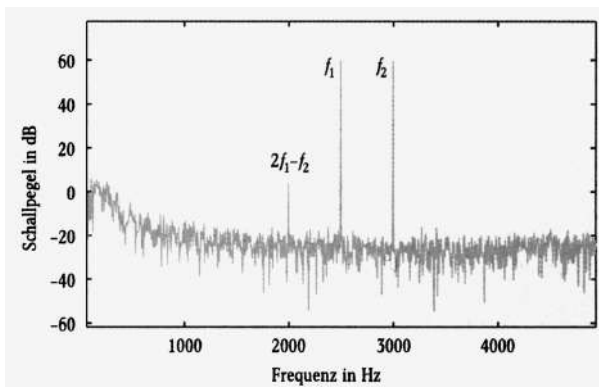


Abbildung 3: Otoakustische Emission. Angeboten werden im abgeschlossenen Gehörgang zwei Sinustöne bei den Frequenzen $f_1 = 2500$ Hz und $f_2 = 3000$ Hz. Bei der Schallaufnahme mit einem empfindlichen Mikrofon erscheint im hier gezeigten gemittelten Spektrum bei gesundem Gehör zusätzlich das kubische Verzerrungsprodukt bei $2f_1 - f_2 = 2000$ Hz. Es erreicht einen Schallpegel von rund 5 Dezibel.

Funktionsstörungen des Außen- und Mittelohres machen sich in einer Schallleitungs-Schwerhörigkeit bemerkbar, die durch HNO-ärztliche Eingriffe oder durch ein einfaches, lineares Hörgerät ausgeglichen werden können. Störungen des Innenohrs oder der nachfolgenden neuronalen Strukturen führen zur Schallempfindungs-Schwerhörigkeit, die wesentlich häufiger vorkommt (ca. 15 % unserer Bevölkerung), aber leider auch schwerer zu kompensieren ist als die Schallleitungs-Schwerhörigkeit (Böhme und Welzl-Müller 1998).

Das Problem liegt in der Veränderung der Nichtlinearität des Innenohrs, die sich auch objektiv mit einem empfindlichen Mikrofon ausmessen lässt. Abbildung 3 zeigt das gemittelte Spektrum des akustischen Signals, das im abgeschlossenen Gehörgang eines normalhörenden Probanden gemessen werden kann, wenn die mit f_1 und f_2 gekennzeichneten reinen Sinustöne als Eingangssignale gegeben werden. Die Verzerrungsprodukte (am ausgeprägtesten bei $2f_1 - f_2$) werden vom Innenohr akustisch abgestrahlt (otoakustische Emissionen) und verschwinden bei Hörstörungen oder bei Unterbrechung der Sauerstoff-Zufuhr des Innenohrs. Die negative Dämpfung bei kleinen Auslenkungen der Basilarmembran und die Ankopplung an äußere Reize lässt sich in guter Näherung durch einen nichtlinearen Van-der-Pol-

Oszillator beschreiben (Uppenkamp et al. 1994). Weil die otoakustischen Emissionen in der Regel nur bei gesunden Ohren auftreten, werden sie auch als objektiver Hörtest bei Neugeborenen eingesetzt. Dazu beschallt man das Ohr mit einem leisen Klick und misst Zeitverzögerung, Frequenzspektrum und Reproduzierbarkeit des akustischen Echos. Otoakustische Emissionen können auch spontan – ohne äußeren Reiz – im abgeschlossenen Gehörgang mit niedrigen Schalldruckpegeln auftreten und weisen keinerlei Beziehung zu dem subjektiv empfundenen Ohrensausen (Tinnitus) auf, das meist mit einer Hörstörung verbunden ist.

Wie funktioniert ein MP3-Player?

Wenn man Musik direkt von der CD über das Internet übertragen will, dauert das ziemlich lange, weil enorme Datenmengen für die genaue Speicherung von Musik benötigt werden. Dabei ist der größte Teil dieser Information gar nicht notwendig, weil unser Ohr nicht merkt, ob sie vorhanden ist oder nicht. Diesen Effekt nutzt die Audio-Kodierung für die Musik-Übertragung über das Internet im MP3-Format aus.

Musik wird in unserem Innenohr in verschiedene Tonhöhenbereiche (Frequenzbänder) zerlegt. Unser Gehirn nimmt die Musikinformation durch die zeitliche Schwankung in diesen Frequenzbändern wahr. Bei einem Musikstück sind diese Frequenzbänder allerdings nicht unabhängig voneinander, sondern schwanken gleichsinnig im Rhythmus der Musik, wobei zumeist ein Frequenzband stärker als die benachbarten Bänder ist. Unser Ohr »hört« daher nur auf dieses stärkere Band und kann die maskierte Information in den benachbarten Frequenzbändern gar nicht wahrnehmen. Dieser Effekt erlaubt dem MP3-Player, die nicht hörbaren Frequenzbänder nicht zu übertragen bzw. lediglich Datenmüll, das so genannte Quantisierungs-Rauschen, unhörbar hineinzupacken. Nur die notwendige Information aus den jeweils dominierenden, vom Ohr auch tatsächlich gehörten Frequenzbändern wird mit hoher Qualität übertragen.

Um nun die Wichtigkeit jedes Frequenzbandes für den Höreindruck festzulegen, benötigt man ein gutes Verständnis der Arbeitsweise des menschlichen Gehörs, ein Hörmodell. Es ist das Herzstück eines jeden MP3-Kodierungsverfahrens und Gegenstand aktueller Forschungsarbeiten. Je besser das Modell und das zugehörige Kodierungsverfahren

ren funktionieren, umso weniger kann das Ohr zwischen dem Original-Signal und dem vom MP3-Player erzeugten Signal unterscheiden, seien es Kastagnettenklänge oder Solo-Piano. Die übertragene Musik ist physikalisch eine ganz andere Musik als das Original, aber sie klingt genauso.

Wie nichtlinear ist das Innenohr?

Dass die Funktion des normalen Innenohrs hochgradig nichtlinear ist, lässt sich in einem einfachen Hörexperiment mit einem Paar Stereolautsprecher demonstrieren (Hörbeispiel unter <http://medi.uni-oldenburg.de/>): Auf einem Lautsprecher wird ein konstanter Sinuston von $f_1 = 2$ kHz abgespielt. Mit dem zweiten Lautsprecher wird dagegen ein Ton mit einer aufsteigenden Frequenz f_2 (*Sweep*) zwischen 2 kHz und 2,8 kHz erzeugt. Werden nun beide Lautsprecher gleichzeitig angeschaltet (lineare Addition des Schallfeldes in der Luft), nimmt man subjektiv (je nach Lautstärke der beiden Stimuli und Funktionszustand des eigenen Mittel- und Innenohres) einen abwärts verlaufenden *Sweep* (kubischer Differenzton $2f_1 - f_2$) und einen Aufwärts-*Sweep* mit niedrigerer Frequenz (quadratischer Differenzton $f_2 - f_1$) wahr.

Diese so genannten Verzerrungsprodukte werden durch eine kompressive Nichtlinearität im Innenohr erzeugt: Für schmalbandige Signale ist die dort angeregte Schwingung proportional zu einer Potenzreihe des Eingangssignals, wobei die dritte Potenz zum kubischen Differenzton führt. Beim normalen Hörvorgang mit natürlichen, breitbandigen Signalen sind die Verzerrungsprodukte unhörbar, weil sie von den stärkeren Komponenten des Eingangs-Schallsignals bei der jeweiligen Frequenz verdeckt (maskiert) werden. Bei der Wahrnehmung arbeitet unser Ohr effektiv linear und vermag, selbst kleinste Nichtlinearitäten im Schallsignal zu erkennen (z.B. den »Klirrfaktor« der HiFi-Anlage kurz vor deren Übersteuerung). Offenbar hat unser Gehirn es gelernt, die Nichtlinearität des Innenohrs und der nachfolgenden Nervenübertragung durch zentrale Signalverarbeitung und kognitive Prozesse auszugleichen. Bei Innenohr-Störungen entfällt die periphere Dynamik-Kompression. Die auf eine funktionierende Kompression eingestellte Signalverarbeitung im Gehirn bleibt aber bestehen, sodass das Gesamtsystem effektiv nichtlinear wird (Derleth et al. 2001).

Als wichtigste Quelle der Nichtlinearitäten im Innenohr und damit der otoakustischen Emissionen gelten die äußeren Haarzellen, die unter Spannungseinfluss zu aktiven Kontraktionen fähig sind, quasi als Resonanzverstärker von Basilarmembran-Schwingungen dienen und so die Empfindlichkeit und die Frequenzauflösung des Hörorgans steigern (Zenner 1994). Die inneren Haarzellen dienen dagegen primär als Mikrofone, die die Bewegung der Basilarmembran in eine Erregung im Hörnerv umsetzen, der diese zum Gehirn weiterleitet. Die genauen mikromechanischen Mechanismen und die Kodierung akustischer Information im Nervensystem sind nach wie vor Gegenstand biophysikalischer Forschung.

Bei dem nachgeschalteten zentralen Hörsystem im Gehirn sind die folgenden Verarbeitungsprinzipien realisiert:

- Frequenzabbildung: Entsprechend der Frequenz-Orts-Transformation auf der Basilarmembran im Innenohr – hohe Frequenzen werden an der Basis der Cochlea abgebildet, niedrige Frequenzen an der Spitze –, werden auf allen Stationen der Hörbahn benachbarte Frequenzen an benachbarten Orten verarbeitet.
- Einhüllenden-Extraktion und Dynamikkompression: In jedem Frequenzband, das zu einem bestimmten Ort auf der Basilarmembran gehört, überträgt der Hörnerv nicht mehr die vollständige Feinstruktur mit der Phaseninformation des Eingangssignals, sondern nur noch seine Einhüllende, d.h. die langsamer schwankende momentane Schallintensität. Dabei wird ein sehr großer Dynamikbereich von bis zu 120 dB Pegeldifferenz zwischen Hörschwelle und Unbehaglichkeitsschwelle erreicht, obwohl jedes einzelne beteiligte Neuron nur maximal etwa 40 dB überdeckt. Die Mechanismen der Adaptation (Anpassen der Empfindlichkeit eines Neurons an den gerade vorliegenden mittleren Eingangspegel) und der Umsetzung der Schallintensität in einen subjektiven Lautheitseindruck sind noch nicht abschließend geklärt.
- Modulationsabbildung: Innerhalb jedes Frequenzbandes wird die Signal-Einhüllende in verschiedene Amplituden-Modulationsfrequenzen aufgespalten. Damit werden über sämtliche Frequenzen hinweg die verschiedenen Rhythmen des Eingangssignals, aber auch periodische Zeitstrukturen (z.B. die Grundfrequenz bei Sprachlauten) systematisch auf benachbarte Nervenzellen im Gehirn abgebildet. Diese zweidimensionale Abbildung (Mittenfrequenz versus Modulationsfrequenz) wurde aufgrund von physiologischen Mes-

- sungen gefunden (Schreiner und Langner 1997) und anhand psychoakustischer Messungen und Modelle bestätigt (Dau et al. 1997).
- Binaurale (räumliche) Abbildung: Um den Ort bzw. die Einfallrichtung einer Schallquelle zu ermitteln, wertet das Gehirn die zwischen beiden Ohren auftretende Laufzeit- und Pegeldifferenz aus und setzt sie in eine innere Karte der akustischen Umwelt um. Dabei wird eine Winkelauflösung von bis zu 1 Grad und eine Zeitgenauigkeit von bis zu 20 Mikrosekunden erreicht – eine der fantastischsten Leistungen unseres Gehörs!

Das Ohr als Spektralapparat oder als Zeit-Analysator

Während frühere Hörtheorien und Erklärungen des Sprachverstehens vorwiegend von einer spektralen Sichtweise des Hörvorgangs ausgingen (z.B. Unterscheidung der Sprach-Vokale durch die Lage der spektralen Maxima/Formanten) und den zeitlichen Aspekt der Hörwahrnehmung als sekundär ansahen, verhält es sich bei modernen Hörtheorien genau umgekehrt: Das Ohr ist nicht nur das schnellste Sinnessystem des Menschen, es kann auch die einem akustischen Signal aufgeprägte zeitliche Information sehr genau verfolgen. So können zeitliche Lücken ab einer Dauer von ca. 5 ms in einem breitbandigen Signal sicher detektiert werden. Begrenzt wird die Zeitauflösung durch die Vor- und Nachverdeckung, d.h. ein Testsignal kann ab ca. 10 ms vor und bis zu 200 ms nach einem (lauteren) Maskierungssignal nicht mehr gehört werden.

Ein Beispiel gegen die rein spektrale Sichtweise ist die Verständlichkeit von *flat-spectrum speech* (Schroeder 1999), d.h. von gefilterter Sprache, deren Kurzzeitspektren ohne spektrale Information, also flach sind (Hörbeispiele unter <http://medi.uni-oldenburg.de>). Ein anderes Beispiel für Sprachverstehen ohne intaktes Sprachspektrum ist bandpassgefilterte Sprache, die in einer spektralen Lücke eines Rauschens dargeboten wird. Umgekehrt wird das Sprachverstehen durch eine Veränderung der Zeitstruktur stärker beeinträchtigt als durch eine spektrale Verformung: So nimmt die Sprachverständlichkeit bei periodisch unterbrochener Sprache ab einer bestimmten Unterbrechungsrate sehr stark ab. Wenn in die zeitlichen Lücken jedoch anstelle der Original-Sprache ein Rauschen mit festem Spektrum gefüllt wird, ist die Zeitstruktur nur noch relativ wenig gestört und unser Gehirn ist in

der Lage, die Sprachinformation wieder zusammenzusetzen (vgl. Hörbeispiel): Das Abwechseln von Sprachsegmenten mit Rauschsegmenten hört sich wie stark verrauschte, kontinuierliche Sprache an, d.h. durch Zufügen (!) von Rauschen wird die Sprache entstört!

Die plausibelste Hörtheorie ist daher eine Kombination der spektralen und zeitlichen Analyse im Sinne einer Demodulation der in jedem Frequenzband vorhandenen Zeit-Information durch das Innenohr mit anschließender Einhüllenden-Analyse in Modulationsfrequenzbändern im Gehirn. Besondere Bedeutung kommt dabei zeitlichen Merkmalen und Modulationsfrequenzen zu, die in mehreren Frequenzbändern gleichzeitig auftreten. Dadurch ist das Ohr in der Lage, akustische Objekte (z.B. Sprache) auch bei sehr ungünstigen Signal-Rausch-Abständen noch sicher zu erkennen, bei denen das mittlere Sprachspektrum schon vollständig verdeckt ist.

Modelle der effektiven Signalverarbeitung

Aus Sicht der Physik möchte man vor allem die Gesetzmäßigkeiten und Mechanismen verstehen, die möglichst vielen der beobachtbaren Phänomene zugrunde liegen, während eine detaillierte Nachbildung jeder einzelnen Nervenzelle zunächst nicht im Vordergrund steht. Ein wichtiges Anliegen der physikalischen Hörforschung ist daher die Entwicklung von quantitativen Hörmodellen, die mit einer möglichst kleinen Zahl von Annahmen und einzustellenden Parametern eine möglichst große Variationsbreite von Experimenten quantitativ vorhersagen und anhand »kritischer« Experimente bestätigt oder falsifiziert werden können. Obwohl dieser Modellansatz damit von vornherein starken Einschränkungen unterliegt, kann er unser derzeitiges Wissen über die effektive Verarbeitung akustischer Information überprüfbar zusammenfassen und zu einem funktionalen Verständnis des Hörvorgangs führen. Der prinzipielle Aufbau eines derartigen Modells der effektiven Verarbeitung im auditorischen System ist in Abbildung 4 dargestellt. Einige der Funktionsblöcke wurden direkt aus physiologischen Erkenntnissen abgeleitet.

So wird z.B. die Funktion der Basilarmembran als effektive Bank von Bandpassfiltern modelliert, die effektive Funktion von Haarzellen und auditorischem Nerv als Kompression und Einhüllendenbildung, die Funktionsprinzipien des Hirnstamms als Modulations-Filterbank

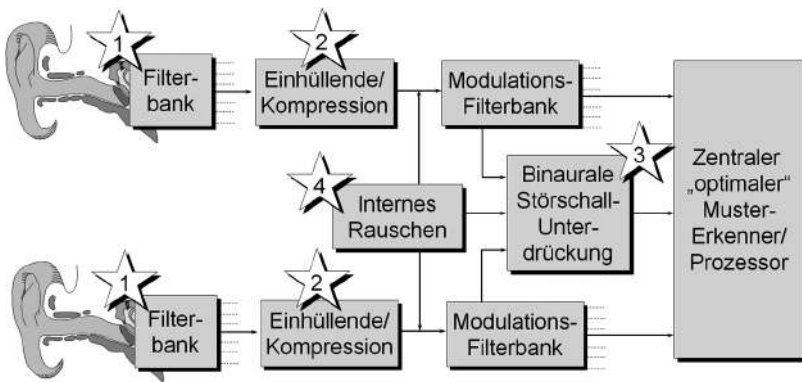


Abbildung 4: Das Ohr aus Sicht des Physikers ist ein Modell der effektiven Verarbeitung im auditorischen System. Es ist durch eine Reihe von parallelen, linearen und nichtlinearen Signalverarbeitungs-Operationen gekennzeichnet, die die Transformation des akustischen Signals in seine interne Repräsentation im Gehirn beschreiben. Diese Repräsentation dient unserem Gehirn als Basis für höhere kognitive Leistungen wie das Verstehen von Sprache, wobei im Modell die Minderung dieser Leistungen allein durch die Unschärfe der internen Repräsentation angesetzt wird. Mögliche Störungen dieser Repräsentation bei Schwerhörigkeit sind durch Sterne gekennzeichnet.

und binaurale Verarbeitung, und schließlich die Funktion des Kortex als Mustererkennung und Interpretation der internen Repräsentation des Schalls. Andere Funktionsblöcke (z.B. das interne Rauschen) und die funktionelle Gestaltung und Parameterwahl sämtlicher Blöcke stammen dagegen aus psychoakustischen Experimenten, zu deren Vorhersage das Modell primär eingesetzt wird. Ebenso stammen die hypothetischen Funktionsstörungen, die bei einem sensorineuralen Hörverlust auftreten können (Sterne in Abb. 4), aus audiologischen Messungen mit schwerhörigen Patienten.

Um den prinzipiellen Rahmen des Modells in ein konkretes, numerisches Verarbeitungsmodell umzusetzen, sind umfangreiche theoretische und experimentelle Arbeiten notwendig. Dabei setzt jede Forschungsgruppe unterschiedliche Schwerpunkte im Detaillierungsgrad einzelner Funktionsblöcke des Modells und bei der Klasse von vorher-sagbaren Experimenten:

- Das in München von Zwicker und Mitarbeitern entwickelte Lautheitsmodell (Zwicker und Fastl 1990) mit seinen Erweiterungen für

- die Beschreibung von Schwerhörigkeit und Schallfluktuationen beschreibt beispielsweise die primär von der Schallintensität, dem Schallspektrum und der Schalldauer abhängige Lautheitswahrnehmung, wobei die Blöcke »Modulationsfilterbank«, »Internes Rauschen« und »binaurale Störschallunterdrückung« aus Abbildung 4 entfallen.
- Das Boston-Modell zur binauralen Informationsverarbeitung von Colburn und Mitarbeitern beschreibt dagegen die verschiedenen Leistungen des binauralen (zweiohrigen) Hörens unter expliziter Modellierung von Neuroneneigenschaften (Colburn 1996). Andere binaurale Funktionsmodelle (z.B. das Bochumer Modell von Blauert und Mitarbeitern (Blauert 1997)) benutzen nachrichtentechnische Funktionselemente, um charakteristische Eigenschaften der binauralen Informationsverarbeitung im Gehirn funktionell zu modellieren.
 - Das Modell der Cambridge-Arbeitsgruppe um Patterson und Meddis setzt dagegen Schwerpunkte bei der Tonhöhenerkennung und dem Übergang von der Wahrnehmung aperiodischer Vorgänge in periodische Vorgänge mit zugehörigem *pitch* (Patterson et al. 1995).
 - Das Oldenburger Perzeptionsmodell legt einen besonderen Schwerpunkt auf die zeitlichen Eigenschaften der Signalverarbeitung im Gehör und bildet eine relativ große Zahl von psychoakustischen Effekten quantitativ nach. Es wurde zunächst am III. Physikalischen Institut in Göttingen und ab 1993 im Graduiertenkolleg Psychoakustik der Universität Oldenburg von Dau und Mitarbeitern (Dau 1997) als Modell der effektiven Signalverarbeitung im auditorischen System entwickelt und auf die Vorhersage von Sprachverständlichkeit bei Normal- und Schwerhörenden (Holube und Kollmeier 1996) sowie das binaurale Hören erweitert. Das Modell basiert auf einer geringen Zahl von Annahmen und Parametern, die in wenigen, »kritischen« Experimenten festgelegt und für die quantitative Beschreibung anderer Experimente nicht mehr variiert werden. Dabei beschreibt das Modell die Transformation des akustischen Signals (z.B. am Eingang zum Innenohr) in eine interne Repräsentation (z.B. ein neuronales Erregungsmuster), auf der dann höhere, kognitive Prozesse aufsetzen, die durch eine »optimale Mustererkennung« (z.B. Vergleich mit gelernten Mustern) quantitativ modelliert werden. Auf diese Weise brauchen das Weltwissen der Versuchsperson und Aspekte wie Aufmerksamkeit und Lernen nicht mehr in

die Modellierung eingebracht zu werden, weil es bei einem »optimalen Detektor« nichts mehr zu verbessern gibt. Die gesamte Imperfektion des Hörvorgangs wird dagegen auf die Nichtlinearität der Signalverarbeitung und das interne, neuronale Rauschen bei der Transformation des akustischen Signals in seine interne Repräsentation reduziert. Dieses Vorgehen hat den entscheidenden Vorteil, dass man sich nur auf denjenigen Teil der Wahrnehmungsleistungen beschränkt, der reproduzierbaren psychophysikalischen Experimenten zugänglich ist und sich zudem möglicherweise physiologisch bestätigen lässt, während die Mehrzahl der komplexen psychischen Einflussfaktoren der menschlichen Wahrnehmung akustischer Ereignisse außen vor bleiben.

Modellierung gestörter Hörfunktionen

Ein möglichst quantitatives Verständnis der gestörten Signalverarbeitung bei Schallempfindungs-Schwerhörigkeit ist sowohl für die Hördiagnostik als auch für die optimale Rehabilitation, z.B. mit »intelligenten Hörgeräten«, unabdingbar. Wegen der Vielzahl der gestörten Einzelleistungen bei Schwerhörigkeit ist es eine besondere Herausforderung, die ursächlichen oder primären Defizite der Hör-Signalverarbeitung von den daraus ableitbaren Defiziten anderer Hörfunktionen zu trennen. Im Rahmen des in Abbildung 4 dargestellten Modellschemas lassen sich nun die vier wichtigsten primären Komponenten von Hörstörungen (Sterne in Abb. 4) wie folgt charakterisieren:

1) Abschwächungswirkung des Hörschadens (lineare Dämpfung): Eine Schallleitungs-Schwerhörigkeit oder ein Ausfall der inneren Haarzellen führt vorwiegend zu einer Sensitivitäts-Verminderung, d.h. einer effektiven Abschwächung des Schalls. Sie kann durch eine entsprechende lineare Verstärkung des Schalls kompensiert werden. Dies bewirkt aber meistens keine zufriedenstellende Wiederherstellung des Hörvermögens, sodass weitere Komponenten betrachtet werden müssen.

2) Kompressionsverlust: Ein Ausfall der äußeren Haarzellen führt zusätzlich zur Abschwächung zu einer Verzerrung: Bei niedrigen Pegeln entfällt die aktive Verstärkung, sodass sich der große Dynamikbereich der akustischen Eingangssignale nicht mehr vollständig im Gehirn abbilden lässt. Dies macht sich beim für die meisten Innen-

ohr-Schwerhörigen typischen *Recruitment*-Phänomen bemerkbar (Hörbeispiel unter <http://medi.uni-oldenburg.de>), bei dem nach dem subjektiven Eindruck »zu leise« bei leichter Erhöhung des Schallpegels bereits der Eindruck »zu laut« folgt. Diese gestörte Lautheitswahrnehmung kann durch eine Multiband-Dynamikkompression in modernen Hörgeräten nur teilweise kompensiert werden, da z.B. die Bandbreiten-Abhängigkeit der unterschiedlichen Lautheitswahrnehmung bei Normal- und Schwerhörigen nicht berücksichtigt wird. Die Entwicklung adäquater Lautheitsmodelle für Schwerhörige und ihre Integration in Hörgeräte ist daher Gegenstand laufender Forschung (Hohmann und Kollmeier 1996).

3) Binauraler Hörverlust: Normalhörende können die an beiden Ohren eintreffenden Signale im Gehirn vergleichen und durch binaurale (beidohrige) Signalverarbeitung den wahrnehmbaren Nachhall verringern und unerwünschte Schalleinfallrichtungen ausblenden. Bei Schwerhörigen kann aber – weitgehend unabhängig von den übrigen bisher genannten Faktoren der Hörstörung – genau diese binaurale Signalverarbeitung gestört sein. Dies bedingt u.a. die eingangs erwähnte Störung des Cocktail-Party-Effektes. In der derzeitigen Routine-Diagnostik und Hörgeräteversorgung mit unabhängigen Geräten auf beiden Seiten wird allerdings dieser Faktor noch nicht berücksichtigt. Erst echt binaurale Hörgeräte versprechen Abhilfe.

4) Zentrale Hörstörung: Selbst bei nur gering gestörter Signalverarbeitung durch das Hörsystem kann bei Schwerhörigen die Auflösung der internen Repräsentation verringert sein, sodass die vom Gehör aufgenommenen und intern repräsentierten Schallsignale nicht mehr adäquat ausgewertet und interpretiert werden können. Dieser Effekt lässt sich ebenso wie andere Unzulänglichkeiten der zentralen Auswerte-Einheit, z.B. geringe Aufmerksamkeit oder mangelndes Training, am ehesten durch ein erhöhtes internes Rauschen modellieren.

Sämtliche der hier genannten Komponenten tragen bei dem individuellen Patienten in unterschiedlichem Maße zu der Hörstörung und beispielsweise zur Verringerung des Sprachverstehens in Störgeräuschen bei. Daher ist es das Ziel aktueller Forschungsarbeiten, effiziente Messmethoden zu entwickeln, mit denen sich jede der Komponenten erfassen lässt, sowie das Modell so anzupassen, dass es das jeweilige Hörvermögen jedes individuellen Patienten korrekt beschreibt (Kießling et al. 1997).

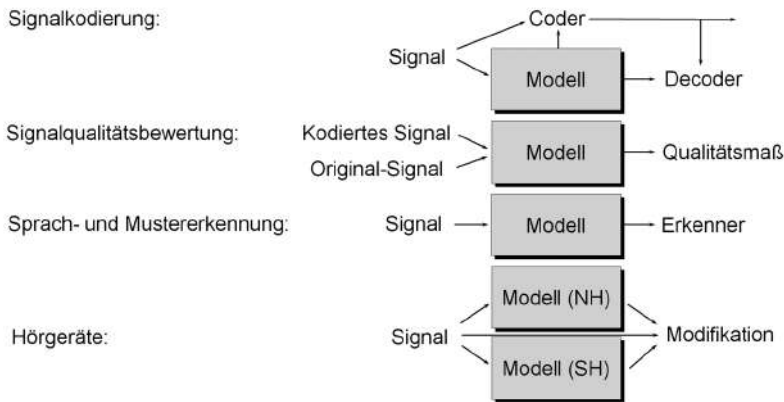


Abb. 5: Anwendungen von Hörmodellen: Durch die quantitative Modellierung des Hörvorgangs können Signale nicht mehr nur auf der akustischen Ebene (Eingang des Hörmodells), sondern auch auf der Wahrnehmungsebene (Ausgang des Modells) beschrieben werden. Dadurch kann die wahrgenommene Qualität einer Signalveränderung objektiv beschrieben (z.B. bei Kodierungen im MP3-Format oder bei der Mobilfunkübertragung) oder die objektive Erkennung relevanter Wahrnehmungsmerkmale erleichtert werden (z.B. bei der künstlichen Spracherkennung). Wenn ein Modell einer Hörstörung vorhanden ist, kann ein modellbasiertes Hörgerät realisiert werden. NH: Normalhörender, SH: Schwerhöriger.

Anwendungen des Perzeptionsmodells

Unter der Voraussetzung, dass das oben beschriebene Perzeptionsmodell eine valide und objektive Beschreibung der Transformation des akustischen Signals in seine interne Repräsentation im menschlichen Gehirn darstellt, erschließen sich eine Reihe von technischen Anwendungen (Abb. 5).

- **Signalkodierung:** Die Bitraten-Reduktion bei der Speicherung von Sprach- und Audiodaten, z.B. mit dem MP3-Verfahren, kodiert das akustische Signal so, dass möglichst geringe Abweichungen zwischen Original und dekodiertem Signal auf der perceptiven Ebene am Ausgang des Hörmodells auftreten, obwohl diese Signale auf der akustischen Ebene am Eingang des Hörmodells sehr unterschiedlich sein können. Bei der standardisierten MP3-Kodierung wird im Wesentlichen das von Zwicker vor mehreren Jahrzehnten entwickelte

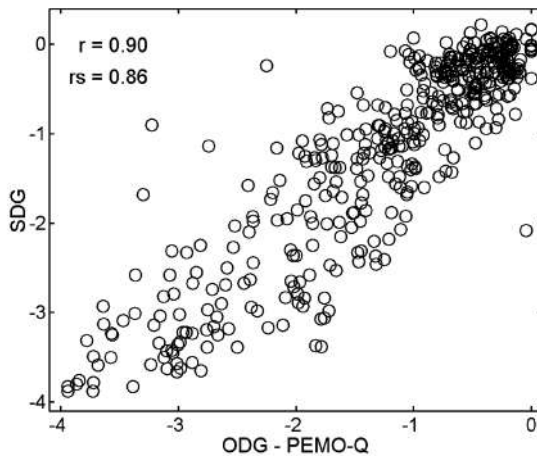


Abbildung 6: Audioqualitäts-Vorhersage: Die Ordinate zeigt das subjektive Qualitätsurteil einer Gruppe normalhörender Probanden (als *subjective difference grade* SDG, d.h. subjektive Verschlechterung in (Schul-)Noten) für eine Test-Datenbank von 433 verschiedenen Audiosignalen, Die Abszisse zeigt die Vorhersage dieser Daten als *objective difference grade* (ODG) auf der Basis des Oldenburger Perzeptionsmodells. Die hohen Korrelations-Indices r und r_s (nahe bei 1) zeigen eine gute mittlere objektive Vorhersage der subjektiven Ergebnisse an (Bildnachweis: Huber und Kollmeier 2006).

Lautheits- und Maskierungsmodell verwendet, um unhörbare Signalbestandteile zu eliminieren und das Quantisierungsrauschen hinter den hörbaren Komponenten zu verstecken.

- Signalqualitäts-Bewertung: Der Unterschied am Ausgang des Hörmodells wird auch bei der objektiven Güte-Beurteilung von kodierten bzw. nichtlinear verarbeiteten Sprach- und Audiosignalen ausgewertet, die bisher nur subjektiv mit aufwändigen Hörexperimenten ermittelt werden konnte. Beispielsweise lässt sich die wahrgenommene akustische Qualität einer Musik-Codierung (z.B. für den MP3-Player oder andere Komprimierungsverfahren von Musik im Internet) durch das Oldenburger Perzeptionsmodell mit dem Computer vorhersagen (objektive Verschlechterung ODG aufgetragen auf der Abszisse in Abbildung 6). Der Vergleich mit der subjektiven Verschlechterung (SDG) auf der Ordinate zeigt im Mittel eine hohe Treffsicherheit der objektiven Vorhersage (Huber und Kollmeier 2006).

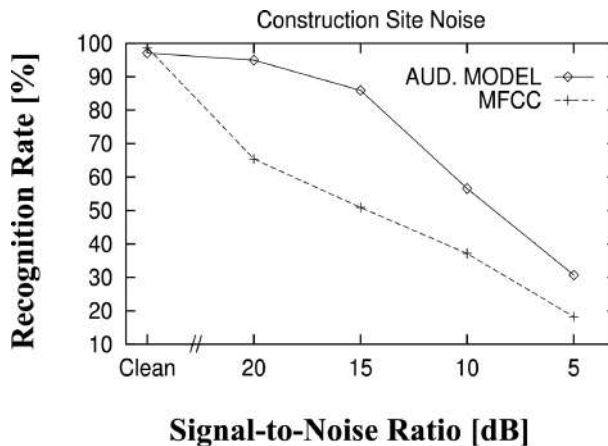


Abbildung 7: Robuste Spracherkennung. Dargestellt ist die Worterkennungsrate eines künstlichen Spracherkenners unter Ruhebedingungen (*Clean*) und als Funktion des Signal-Rausch-Abstandes für Sprache in Baustellenlärm (*Construction Site Noise*) für die konventionelle Sprachvorverarbeitung (MFCC, gestrichelte Linie) und für das Oldenburger Perzeptionsmodell (AUD. MODEL, durchgezogene Linie) (Bildnachweis: Tchorz und Kollmeier 1999).

- Sprach- und Mustererkennung: Dem Computer werden »Ohren verliehen«, indem nicht eine technische Darstellung des Sprachsignals, sondern die interne Repräsentation vom Ausgang des Gehörmodells als Eingangssignal für einen Spracherkennungs- bzw. Mustererkennungs-Algorithmus benutzt wird. Die Robustheit unseres Ohres gegenüber Störgeräuschen und Änderungen der Raumakustik soll damit auf den Computer übertragen werden: Abbildung 7 zeigt die Worterkennungsrate in Prozent als Funktion des Signal-Rausch-Abstandes. Bei der konventionellen Sprachvorverarbeitung mit dem so genannten MFCC-Verfahren (*Mel Frequency Cepstral Coefficients*, gestrichelte Linie) sinkt die Erkennungsrate mit zunehmendem Rauschen viel eher als bei der Sprachvorverarbeitung mit dem Oldenburger Perzeptionsmodell (durchgezogene Linie, aus: Tchorz und Kollmeier 1999). Bei einer stark eingeschränkten Zahl von möglichen Alternativen des zu erkennenden Wortes oder Satzes erreicht der Computer sogar die Erkennungsleistung des menschlichen Gehörs.

- »Intelligente Hörgeräte«: Ziel der Signal-Manipulationen im Hörgerät sollte es sein, den Unterschied der internen Repräsentation am Ausgang des Hörmodells zwischen Normal- und Schwerhörigen zu minimieren. Dies setzt jedoch neben dem Modell für das normale Gehör auch ein Modell des gestörten Hörvermögens voraus, das individuell angepasst wird. Obwohl inzwischen digitale Hörgeräte mit einem vergleichbaren Konzept auf dem Markt sind, ist dies noch ein Bereich aktueller Forschungs- und Entwicklungsarbeiten.

Störgeräuschunterdrückung in Hörgeräten der Zukunft

Ein wichtiges Beispiel für die Umsetzung von Hörmodellen in die Praxis ist die Störgeräuschunterdrückung in Hörgeräten, die möglichst zu einer Wiederherstellung des Cocktail-Party-Effektes bei Schwerhörigen führen und die Auswirkungen der reduzierten binauralen Interaktion und des erhöhten internen Rauschens zumindest teilweise kompensieren soll. Obwohl eine für alle möglichen akustischen Störschall-Nutzschall-Konfigurationen wirksame Störunterdrückung auch für andere Anwendungen in der Sprachkommunikation (z.B. automatische Spracherkennung) sehr erstrebenswert ist, zählt dies zu den großen noch ungelösten Problemen der Akustik.

Einen vielversprechenden Ansatz bietet die Analyse mit dem so genannten Amplituden-Modulations-Spektrogramm (AMS), das in jedem Frequenzband die zeitlichen Fluktuationen in verschiedene Modulations-Frequenzen zerlegt. Es entspricht der bereits in Abbildung 4 dargestellten Modellvorstellung, dass in jedem auditorischen Frequenzband die zeitliche Einhüllenden-Struktur durch eine Modulations-Filterbank ausgewertet wird. In dieser Darstellung ist Sprache durch in mehreren Frequenzbändern vorhandene kohärente Modulationen im Modulationsfrequenzbereich von 4 Hz (Silbenfrequenz) und mehreren hundert Hz (Sprachgrundfrequenz mit Harmonischen) gekennzeichnet. Störgeräusche weisen hingegen in der Regel weniger kohärente Modulationen und auch ein anderes Modulationsspektrum auf, sodass sich diese Unterschiede gut für die Störgeräuschunterdrückung ausnutzen lassen. Der Vorteil dieses Verfahrens ist seine Anwendbarkeit auch für monaurale (einkanalige) Mikrofonsignale und für fluktuierende Hintergrundgeräusche (Hörbeispiel unter <http://www.medi.uni-oldenburg.de>). Der Nachteil des Verfahrens ist jedoch

der hohe Rechenaufwand. Außerdem versagt es, wenn das Hintergrundgeräusch selbst Sprache (z.B. ein weiterer Sprecher) ist.

Daher bietet ein binaurales Verfahren Vorteile, bei dem die Signale an beiden Ohren aufgenommen werden und in einer zentralen Recheneinheit so gefiltert werden, dass die von vorn kommenden Signale, d.h. der Nutzschaall, verstärkt und der von anderen Richtungen kommende Störschaall unterdrückt wird (Wittkop et al. 1997). Damit wird das binaurale Hören gewissermaßen simuliert. Ein derartiges binaurales Hörgerät hat deutliche Vorteile gegenüber zwei unabhängigen Hörgeräten auf beiden Seiten und erst recht gegenüber einem monauralen Hörgerät. Ein binaurales Hörgerät ist inzwischen mit eingeschränkter Funktionalität kommerziell erhältlich. Die Universität Oldenburg hat als Vorstudie für derartige Geräte zusammen mit insgesamt 29 Partnern des EU-Projekts HEARCOM einen wesentlich leistungsfähigeren Prototypen entwickelt (Abb. 8 und Hörbeispiel für die Verarbeitungsleistung binauraler Hörgeräte-Algorithmen unter <http://medi.uni-oldenburg.de>). Für ein kommerzielles binaurales Hörgerät ist jedoch eine drahtlose Verbindung zwischen den Geräten nötig, die wegen der erforderlichen Stromaufnahme bisher nur mit großen Einschränkungen realisiert werden konnte.

Wie wird die Entwicklung der Hörgeräte weitergehen? Es ist abzusehen, dass das Hörgerät nur noch eine Option eines *Personal Communication Device* der Zukunft darstellen wird, in dem Mensch-Maschine- und Mensch-Mensch-Kommunikationsfunktionen wie MP3-Player, Mobiltelefon, Laptop und eben Hörgerät in einem tragbaren und durch Sprachsteuerung bedienbaren Gerät verschmelzen. Zukünftige Hörgeräte werden zudem binaural sein und neben der modellgesteuerten Dynamikkompression eine an die jeweilige Kommunikationssituation optimal angepasste Stör-Reduktion mit automatischer Programmwahl aufweisen. Auch Cochlea-Implantate, »Hörprothesen« für beidseits ertaubte Patienten, die durch direkte Stimulation des Hörnervs eine akustische Wahrnehmung ermöglichen, werden davon profitieren.

Einen wesentlichen Beitrag für diese Entwicklung leistet der Forschungs- und Entwicklungscluster *Auditory Valley* – Hören in Niedersachsen, der die Expertisen in der Region Oldenburg/Hannover zu Hörsystemen und Hördiagnostik zusammenführt. Mit dem Kompetenzzentrum für Hörgeräte-Systemtechnik HörTech, der Medizinischen Hochschule Hannover, der Abteilung Medizinische Physik der



Abbildung 8: Digitales, binaurales Master-Hörgerät, das im BMBF-geförderten Kompetenzzentrum HörTech für die Entwicklung verschiedener digitaler Hörgerät-Algorithmen genutzt wurde. Im EU-Projekt Hearcom dient es jetzt 29 Partnern als gemeinsame Hard- und Software-Plattform zum Vergleich unterschiedlicher Algorithmen (Bildnachweis: <http://www.hearcom.eu>).

Universität Oldenburg sowie den Hörzentren Hannover und Oldenburg ist die Region das bedeutendste Zentrum für angewandte Hörforschung in Europa. Der Verbund umfasst die Wertschöpfungskette hochtechnologischer Hörsysteme von der Forschung über die Entwicklung bis hin zur Produkteinführung, die Versorgung der Betroffenen sowie Evaluation und Qualitätssicherung. In einem vorwettbewerblichen Wissenstransfer werden gemeinsam mit den Partnern aus Forschung und Industrie die Möglichkeiten und Trends in der Hörgeräte-Versorgung ausgelotet. Kooperationspartner sind Industrieunternehmen, die insgesamt 92 Prozent des Weltmarktes der Hörgeräte und 65 Prozent des Weltmarktes für Cochlea-Implantate abdecken. Die jüngste Ansiedlung der Fraunhofer-Projektgruppe Hör-, Sprach- und Audiotechnologie erlaubt zudem die Einbindung von Techniken aus den Bereichen MP3, Bluetooth, etc. Ein Hörgarten ist für die Öffentlichkeit eingerichtet und ermöglicht einen spielerischen Zugang zum Thema (Abb. 9). Ein weiteres Angebot für die Öffentlichkeit ist ein



Abbildung 9: Die Welt des Hörens im Hörgarten. Am Oldenburger Haus des Hörens weckt eine Art akustischer Themenpark mit zahlreichen Exponaten das Hörbewusstsein der Öffentlichkeit.

Rechts: Zu den nicht nur akustisch interessanten, sondern auch ästhetisch ansprechenden Installationen zählen runde Helmholtz-Resonatoren (hier mit dem Autor), die je nach Größe bestimmte Geräusche aus der Umgebung herausfiltern und diese verstärken. Bei den großen Kugeln und Röhren werden die tiefen Töne gefiltert und verstärkt, bei den kleinen die hohen Töne, wobei die Frequenzen den harmonischen Teiltönen bei einem Instrumentenklang entsprechen. Ein weiteres Exponat ist der Flüsterverspiegel, bestehend aus zwei großen, im Abstand von fast 40 Metern installierten akustischen Hohlspiegeln, die es ermöglichen, Schall über weite Entfernungen hinweg gezielt wahrnehmen zu können. Die Reflektoren bündeln den Schall, so dass man sich von dem einen Brennpunkt vor dem ersten Hohlspiegel aus im Flüsterton mit einer zweiten Person unterhalten kann, deren Kopf sich in der Nähe des zweiten Brennpunkts vor dem zweiten Hohlspiegel befindet. Ein Modell des Innenohrs und eine Mittelohrpauke veranschaulichen Funktionsweisen des Hörorgans, ebenso der binaurale Teich (binaural = beidohrig), an dem sich zeigen lässt, dass man zwei Ohren zur Ortung und Identifizierung von akustischen Objekten im Raum benötigt. Der Hörthron verstärkt alle Geräusche der Umgebung um 14 Dezibel. Der Schall wird über zwei große Trichter (Durchmesser jeweils 1 m) und eine Röhre an die Ohren der Person geführt, die zwischen den Hörrohren Platz nimmt. Durch die Trichterwirkung wird nicht nur der ankommende Schall verstärkt, sondern durch den großen Abstand zwischen den beiden Hörtrichtern kann man auch die Richtung von Geräuschen viel genauer bestimmen, weil der zwischen beiden Ohren auftretende Laufzeitunterschied für eine bestimmte Einfallsrichtung bis zu 10-fach vergrößert wird. An einer original englischen Telefonzelle können sich die Besucher des Hörgartens über die Exponate sowie Hören, Schwerhörigkeit und Lärm informieren (Bildnachweis und weitere Informationen: <http://www.hoergarten.de>).

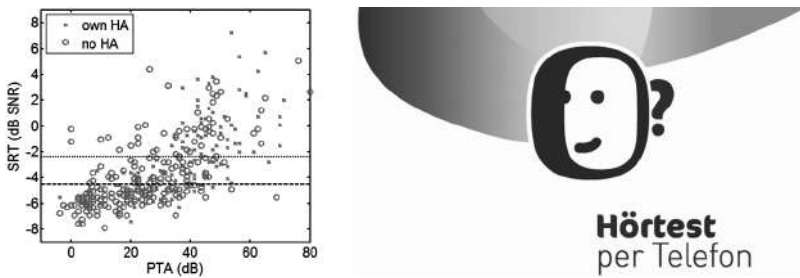


Abbildung 10: Hörtest per Telefon. Das Diagramm zeigt den Zusammenhang zwischen dem Sprachhörvermögen im Störschall, der mit dem Hörtest per Telefon ermittelt wird (SRT, Ordinate) und dem mittleren Tongehör in Ruhe (PTA, Abszisse). Die Daten stammen von insgesamt 326 Patienten, Hörgeräte-Trägern (x) und unversorgten Patienten (o), deren Tongehör in Ruhe (Audio-gramme) am Hörzentrum Oldenburg erhoben wurden und die anschließend freiwillig am Telefontest teilnahmen. Die gestrichelten Linien zeigen die Beurteilungsgrenzen für den Telefontest. Im Mittel hängen Sprachhörvermögen im Störschall und das Tongehör in Ruhe eng zusammen. Manche im Ton-audiogramm unauffällige Patienten zeigen aber im Störschall beträchtliche Defizite und umgekehrt.

Hörtest per Telefon (Abb. 10). Dessen Grundlagen hat das EU-Projekt HEARCOM entwickelt. Für Deutschland hat das Kompetenzzentrum HörTech den Test angepasst. Bereits über 35 000 Anrufer seit der Freischaltung im Juli 2008 ließen ihr Hörvermögen schnell und anonym untersuchen. Der Anruf bei der Hörtestnummer dauert nur wenige Minuten, dann erhält der Anrufer eine wissenschaftlich fundierte Einschätzung, wie es um sein Gehör bestellt ist. Das Besondere und Neuartige an dem Verfahren ist, dass es das Sprachverstehen unter Störgeräuschen überprüft und damit genau die alltäglichen Gesprächssituationen simuliert, bei denen Hörstörungen am häufigsten auftreten. Der Anrufer muss Sprache in einer Geräuschkulisse verstehen. Ihm werden Ziffern vor einem mehr oder weniger starken Störgeräusch eingespielt und er muss die Lösungen über die Zifferntastatur seines Telefons eingeben. Bei richtigen Eingaben wird der Schwierigkeitsgrad erhöht.

Ein Pionier der Hör- und Sprachkommunikations-Forschung

Der Göttinger Physik-Professors Manfred R. Schroeder (1926-2009) leistete mit seiner Forschungs- und Lehrtätigkeit einen maßgeblichen Beitrag zur Entwicklung der Hörgeräte der Zukunft und anderer »intelligenter Audiosysteme«, denn seine Absolventen und ehemaligen Mitarbeiterinnen und Mitarbeiter zählen heute zum personellen Kern des *Auditory Valley*. Manfred Schroeder war maßgeblich an einer Reihe bahnbrechender Publikationen und Erfindungen in den Gebieten raumakustische Messverfahren, Sprachkodierung, physiologische und psychologische Akustik und Computergrafik beteiligt. Beispielsweise sind seine Erfindungen zum künstlichen Nachhall in modernen digitalen Hallgeräten enthalten und eine Weiterentwicklung seiner gemeinsam mit Bishnu Atal entwickelten Methode des *Linear Predictive Coding* ist in jedem modernen Mobiltelefon zu finden. In der Hörforschung sind die Schroeder-Phasen ein feststehender Begriff. Mit ihnen wird in aktuellen Forschungsarbeiten die Funktion der Basilar-membran psychoakustisch und physiologisch untersucht. Zudem hat Manfred Schroeder den Transfer von physikalischen Ideen in erfolgreiche Produkte der Industrie vorgelebt, indem er in der vorlesungsfreien Zeit bei den *Bell Telephone Laboratories* in Murray Hill, USA beschäftigt war und während des Semesters in Göttingen lebte. Seinem Andenken ist dieser Beitrag gewidmet.

Die beschriebene Forschung und Entwicklungen wurden gefördert von der DFG, dem BMBF, der EU und dem Land Niedersachsen. Herzlicher Dank allen Mitarbeiterinnen und Mitarbeitern der Medizinischen Physik, der Hörzentrum Oldenburg GmbH, dem Kompetenzzentrum HörTech, der Fraunhofer-Projektgruppe für Hör-, Sprach- und Audiotechnologie und allen Partner-Institutionen im *Auditory Valley*.

Literatur

- Blauert J. (1997) *Spatial Hearing*. Cambridge (MIT Press).
- Böhme G., Welzl-Müller K. (1998) *Audiometrie. Hörprüfungen im Erwachsenen- und Kindesalter* Bern (Verlag Hans Huber).
- Colburn H.S. (1996) *Binaural Models*. In: *Auditory Computation*, hrsg. von H.L. Hawkins, T.A. McMullen, A.N. Popper, R.R. Fay. New York (Springer), S. 332-400.
- Dau T., Kollmeier B., Kohlrausch A. (1997) Modeling auditory processing of amplitude modulation: I. Detection and masking with narrow band carrier. *Journal of the Acoustical Society of America* 102, S. 2892.
- Derleth R.-P., Dau T., Kollmeier B. (2001) Modeling temporal and compressive properties of the normal and impaired auditory system. *Hearing Research* 159, S. 132.
- Huber R., Kollmeier B. (2006) *IEEE Transactions on Audio, Speech and Language* 14, S. 1902.
- Hohmann V., Kollmeier B. (1996) Perceptual models for hearing aid algorithms. In: *Psychoacoustics, Speech and Hearing Aids*, hrsg. von B. Kollmeier. Singapore (World Scientific), S. 193-202.
- Holube I., Kollmeier B. (1996) Speech intelligibility prediction in hearing-impaired listeners based on a psychoacoustically motivated perception model. *Journal of the Acoustical Society of America* 100, S. 1703.
- Kießling J., Kollmeier B., Diller G. (1997) *Versorgung und Rehabilitation mit Hörgeräten*. Stuttgart (Thieme Verlag).
- Patterson R.D., Allerhand M., Giguere C. (1995) Time-domain modelling of peripheral auditory processing: A modular architecture and a software platform. *Journal of the Acoustical Society of America* 98, S. 1890.
- Schreiner C.E., Langner G. (1997) Laminar fine structure of frequency organization in auditory midbrain. *Nature* 388, 383.
- Schroeder M.R. (1999) *Computer Speech: Recognition, Compression, Synthesis*. Berlin (Springer).
- Uppenkamp S., Neumann J., Kollmeier B. (1994) Chirp Evoked Otoacoustic Emissions. *Hearing Research* 78, S. 210.
- Tchorz J., Kollmeier B. (1999) A model of auditory perception as front end for automatic speech recognition. *Journal of the Acoustical Society of America* 106, S. 2040.
- v. Helmholtz H. (1870) *Die Lehre von den Tonempfindungen als physiologische Grundlage für die Theorie der Musik*. Braunschweig (Vieweg).
- Wittkop T., Albani S., Hohmann V., Peissig J., Woods W.S., Kollmeier B. (1997) *Speech Processing for Hearing Aids: Noise Reduction Motivated by Models of Binaural Interaction*. *Acustica united with Acta acustica* 83(4), S. 684.
- Zenner H. (1994) *Hören*. Stuttgart (Thieme Verlag).
- Zwicker E., Fastl H. (1990) *Psychoacoustics – Facts and Models*. Berlin (Springer).

Die Autoren

Eckart Altenmüller, geboren 1955 in Rottweil am Neckar, promovierte nach dem Medizinstudium in Tübingen, Paris und Freiburg im Breisgau und dem Musikstudium in Freiburg (Hauptfach Querflöte) 1983 über die Gangentwicklung bei Kleinkindern. Während der Assistenzzeit in der Abteilung für klinische Neurophysiologie in Freiburg entstanden die ersten Arbeiten zur Hirnaktivierung beim Musikhören. Von 1985 bis 1994 absolvierte Eckart Altenmüller an der Universität Tübingen die Facharztzeit für Neurologie und habilitierte sich 1992 im Fach Neurologie. Seit seiner Berufung als Direktor des Instituts für Musikphysiologie und Musiker-Medizin der Hochschule für Musik und Theater Hannover im Jahr 1994 sind zahlreiche Arbeiten zum auditiven und sensomotorischen Lernen, zur Störung der Musikverarbeitung nach Schlaganfällen und zur emotionalen Verarbeitung von Musik entstanden. Seit 2003 sind die neuropsychologischen Grundlagen der Gestaltung und Wahrnehmung von starken Emotionen in der Musik ein weiteres wichtiges Forschungsthema.

Eckart Altenmüller ist Mitglied zahlreicher nationaler und internationaler Gremien. 2005 wurde er zum Mitglied der Göttinger Akademie der Wissenschaften ernannt und zum Präsidenten der Deutschen Gesellschaft für Musikphysiologie und Musiker-Medizin gewählt.

Aaron Ciechanover, geboren 1947 in Haifa, studierte Medizin in Jerusalem, dann aber entschied er sich für die Grundlagenforschung und promovierte – unterbrochen von klinischen Praktika und vom dreijährigen Wehrdienst – 1981 im Fach Biochemie am Technion (*Israel Institute of Technology*) in Haifa. In dieser Zeit gelangen die maßgeblichen Schritte der Entdeckung und Charakterisierung des Ubiquitin-Systems. Als Postdoktorand am *Massachusetts Institute of Technology MIT*, Cambridge, USA, war er an der Aufklärung rezeptorvermittelter Endocytose beteiligt. Seine eigene Arbeitsgruppe baute er ab 1984 in Israel in der Fakultät für Medizin am Technion auf, wurde 1992 ordentlicher Professor und widmet sich wieder dem Gebiet der Ubiquitin-Forschung, das zwischenzeitlich die Aufmerksamkeit vieler Arbeitsgruppen auf sich gezogen hatte. Für seine Erkenntnisse zum intrazellulären Proteinabbau erhielt er 2004 zusammen mit seinen

Mentoren Avram Hershko und Irwin Rose den Nobelpreis für Chemie. Als Gastprofessor forschte und lehrte Aaron Ciechanover an namhaften medizinischen Instituten, er ist Mitglied zahlreicher Akademien auf der ganzen Welt und hält über 20 Ehrendoktorwürden und Honorarprofessuren.

Ulrich Christensen, geboren 1954 in Peine, studierte Physik und promovierte 1980 in Braunschweig. Als Postdoktorand war er an der *Arizona State University*, USA und an der Universität Mainz, wo er sich 1985 im Fachbereich Geophysik habilitierte. Mit einem Heisenberg-Forschungsstipendium wechselte er in Mainz ans Max-Planck-Institut (MPI). 1992 erhielt er einen Ruf an die Universität Göttingen, wo er als Professor für Geophysik mit dem Gottfried-Wilhelm-Leibniz-Preis der Deutschen Forschungsgemeinschaft ausgezeichnet wurde. Seit 2002 ist er Direktor am MPI für Sonnensystemforschung in Katlenburg-Lindau, davon 2005-2007 geschäftsführender Direktor. Sein Forschungsgebiet ist der Aufbau und die Dynamik des Inneren der Planeten (einschließlich der Erde). So untersucht er beispielsweise Konvektionsströmungen im Silikatmantel der Erde, die Erzeugung planetarer Magnetfelder in einem Dynamoprozeß und die Strömungen in den großen Gasplaneten mit numerischen Simulationen und schließt auf den inneren Aufbau der Planeten aus geodätischen und seismologischen Daten.

Ulrich Christensen ist Mitglied der Göttinger Akademie der Wissenschaften, der Deutschen Akademie Leopoldina und Ehrenmitglied der *European Union of Geoscience*.

Ian Thomas Baldwin, Direktor und Wissenschaftliches Mitglied am Max-Planck-Institut für chemische Ökologie und Honorarprofessor an der Universität Jena, ist Pionier bei der Beantwortung ökologischer Fragestellungen mittels praktischer Anwendung moderner Pflanzen genomforschung und Methoden der grünen Gentechnik.

Ulf Diederichsen, geboren 1963 in München, studierte Chemie in Freiburg im Breisgau, wurde 1993 an der Eidgenössischen Technischen Hochschule (ETH) Zürich promoviert und ging als Postdoktorand an die Universität Pittsburgh in Pittsburgh, USA. Er habilitierte sich 1999 an der Technischen Universität München mit der Habilitationsschrift »Lineare Nucleinsäure-Analoga mit peptidischem Rückgrat«. Zu-

nächst vertrat er eine Professur an der Ludwig-Maximilians-Universität München, wechselte anschließend an die Universität Würzburg und ist seit 2001 Professor am Institut für Organische Chemie der Universität Göttingen. Sein großer Kreis von Mitarbeitern beschäftigt sich mit der Synthese und Modifikation von Biooligomeren, mit den molekularen Vorgängen der DNA-Bindung, der Modifikation von Proteinen, den Protein- bzw. Peptid-Lipid-Interaktionen in Membranen und mit der Synthese von Nucleotidderivaten.

Ulf Diederichsen wurde ausgezeichnet mit dem Preis der Hellmut-Bredereck-Stiftung (1999) und einem Forschungsstipendium der Karl Winnacker-Stiftung (1999-2000) und verbrachte ein halbes Jahr als *Goering Visiting Professor* an der *University of Wisconsin* in Madison, USA.

Julia Fischer, geboren 1966 in München, studierte Biologie in Berlin und Glasgow und wurde 1996 an der Freien Universität Berlin promoviert. Während ihrer Zeit als Postdoktorandin an der *University of Pennsylvania*, Philadelphia, USA betrieb sie 18 Monate Feldforschung an freilebenden Pavianen in Botswana. Ein Habilitationsstipendium führte sie 2000 an das Max-Planck-Institut für evolutionäre Anthropologie in Leipzig, wo sie sich 2004 an der Universität habilitierte. Im gleichen Jahr erhielt sie ein Heisenbergstipendium der DFG sowie einen Ruf an das Deutsche Primatenzentrum und an die Universität Göttingen. Ihr Forschungsschwerpunkt ist die Evolution kommunikativer und kognitiver Fähigkeiten bei Primaten. Sie war Mitglied der Jungen Akademie und wurde 2007 in die Berlin-Brandenburgische Akademie der Wissenschaften aufgenommen. Als Mitglied im Hochschulrat der Ludwig-Maximilians-Universität München und im Beirat »Frauen in der Wissenschaft« der Robert-Bosch-Stiftung engagiert sie sich auch in der Wissenschaftspolitik.

Jens Frahm, geboren 1951 in Oldenburg, studierte Physik in Göttingen und wurde 1977 im Fachbereich physikalische Chemie promoviert. Als Wissenschaftlicher Mitarbeiter am Göttinger Max-Planck-Institut für biophysikalische Chemie (MPIIbpc), leitete er 1982 bis 1992 eine eigene Arbeitsgruppe. Zwei große Geldgeber erlaubten ihm in dieser Zeit, seine methodische Grundlagenforschung für die Magnetresonanztomografie bei Mensch und Tier voranzutreiben und ein dafür geeignetes eigenes Institut auf dem Gelände des Max-Planck-

Instituts zu bauen: das Bundesministerium für Forschung und Technologie (jetzt BMBF) und die Volkswagen-Stiftung. Seit 1993 leitet Jens Frahm die Biomedizinische NMR Forschungs GmbH, eine gemeinnützige, von der Max-Planck-Gesellschaft getragene Einrichtung, die sich weitgehend aus den Lizenzen für eigene Patente finanziert. Er habilitierte sich 1994 an der Georg-August-Universität Göttingen, wo er seit 1997 als außerordentlicher Professor lehrt.

Jens Frahm wurde unter anderem mit der Goldmedaille der Internationalen *Society for Magnetic Resonance in Medicine* (1991), dem Karl Heinz Beckurts-Preis (1993), dem Niedersächsischen Staatspreis (1996) und für seine Beiträge zur Multiple-Sklerose-Forschung mit dem Werner Sobek-Preis (2005) ausgezeichnet. Außerdem wurde er 2005 in die Akademie der Wissenschaften zu Göttingen berufen.

Markus Hartl, geboren 1978 in Wien, studierte Biologie an der Universität Wien. Als Doktorand und wissenschaftlicher Mitarbeiter am Max-Planck-Institut für chemische Ökologie in Jena in der Arbeitsgruppe von Professor Ian T. Baldwin untersuchte er in Labor- und Freilandversuchen in Deutschland und den USA die genetische Grundlage von pflanzlichen Abwehrmechanismen gegen Schadinsekten. 2005 erhielt er auf dem 3. Symposium der *International Max Planck Research School* den Preis für die beste Präsentation. 2010 wurde er in Jena promoviert. Seitdem ist er Postdoktorand im Emmy Noether-Projekt des Arbeitskreises von Dr. Iris Finkemeier an der Ludwig-Maximilians-Universität in München. Dort erforscht er den Einfluss post-translationaler Veränderungen auf den pflanzlichen Stoffwechsel und dessen Signalmetaboliten.

Jan-Wolfhard Kellmann, promovierter und habilitierter Biologe, ist seit 2004 Forschungskoordinator am Max-Planck-Institut für chemische Ökologie in Jena, wo er mit Markus Hartl zusammenarbeitete. Seine Forschungsgebiete sind die Biochemie und Pathologie von Pflanzen, insbesondere Viruserkrankungen. Schon seit Beginn der Freisetzungen transgener Pflanzen zu Forschungszwecken in Deutschland Anfang der 1990er Jahre ist er in die Thematik der grünen Gentechnik involviert.

Birger Kollmeier, geboren 1958 in Minden, studierte gleichzeitig Physik und Medizin in Göttingen und promovierte 1986 bei Man-

fred R. Schroeder im Fach Physik über Psychoakustik und 1989 im Fach Medizin über Hörgeräte, nicht ohne als Fulbright-Stipendiat in St. Louis/USA Auslandserfahrung gesammelt zu haben. 1991 habilitierte er sich im Fach Physik mit dem Thema Messmethodik, Modellierung und Verbesserung der Verständlichkeit von Sprache. Seit dem Sommersemester 1993 hat er die Professur für Angewandte Physik/Experimentalphysik am Fachbereich Physik der Universität Oldenburg inne und ist Leiter der Abteilung Medizinische Physik, Wissenschaftlicher Leiter der HörTech gGmbH und der Fraunhofer-Projektgruppe für Hör-, Sprach- und Audiotechnologie sowie Vorstandsmitglied der Deutschen Gesellschaft für Audiologie (DGA). Seine Forschungsschwerpunkte sind Sprachperzeption, Hörgeräte-Systemtechnik, und medizinisch-physikalische Hör- und Hirn-Diagnostik.

Norbert Krupp, promovierter Physiker, ist wissenschaftlicher Mitarbeiter in der Abteilung Planeten und Kometen am Max-Planck-Institut für Sonnensystemforschung in Katlenburg-Lindau. Er ist an der Auswertung diverser internationaler Weltraummissionen beteiligt und für die Öffentlichkeitsarbeit des Instituts zuständig.

Christoph Marksches, geboren 1962 in Berlin, studierte evangelische Theologie, klassische Philologie und Philosophie in Marburg, Jerusalem, München und Tübingen. Er wurde 1991 im Fach Theologie promoviert, habilitierte sich 1994 und wurde im selben Jahr zum evangelischen Pfarrer ordiniert. 1995 wurde er als Professor für Kirchengeschichte an die Universität Jena berufen. Nach einer Station in Heidelberg ist er seit 2004 Professor für Ältere Kirchengeschichte (Patristik) an der Humboldt-Universität zu Berlin, als deren Präsident er 2006-2010 wirkte. Forschungsaufenthalte führten ihn nach Jerusalem, Oxford und Princeton.

Christoph Marksches arbeitet in diversen nationalen und internationalen Gremien in der Wissenschaftsförderung, Kirche und Politik mit, z.B. im *Advisory Board* des *Israel Democracy Institute*. Er ist Mitglied der Berlin-Brandenburgischen Akademie der Wissenschaften, weiterer Akademien und Institute. 2001 erhielt er den Gottfried-Wilhelm-Leibniz-Preis der Deutschen Forschungsgemeinschaft, 2007 die Ehrendoktorwürde der Fakultät für Orthodoxe Theologie der Lucian-Blaga-Universität in Sibiu, Rumänien, 2010 den Theologischen Preis

der Salzburger Hochschulwochen und 2011 die Ehrendoktorwürde der Theologischen Fakultät der Universität Oslo.

Eva-Maria Neher, geboren 1950 in Mülheim an der Ruhr, studierte Mikrobiologie und Biochemie an der Universität Göttingen. Nach ihrer Promotion erhielt sie ein Post-doc-Stipendium am Max-Planck-Institut für biophysikalische Chemie und wechselte danach an die Medizinische Fakultät der Universität Göttingen. Nach einer mehrjährigen Familienpause erteilte Eva-Maria Neher an der Freien Waldorfschule Göttingen Experimentalunterricht in Biologie und Chemie und entwickelte das wissenschaftlich-didaktische Konzept zur EXPO-Ausstellung »Faszination Pflanzenzüchtung«. Sie ist Gründerin des XLAB, Experimentallabor für junge Leute, dessen Leitung und Geschäftsführung sie seit 2000 innehat.

Für ihr Engagement für die naturwissenschaftliche Bildung wurde ihr 2007 der Niedersächsische Staatspreis verliehen. 2009 erhielt sie eine Honorarprofessur der Fakultät für Chemie an der Universität Göttingen.

Hans-Peter Nollert, geboren 1957 in Karlsruhe, studierte Physik und promovierte 1990 bei Professor Hanns Ruder in Tübingen. Er habilitierte sich 2000 in Tübingen und hält dort Vorlesungen zu Relativitätstheorie, relativistischer Astrophysik und wissenschaftlicher Visualisierung und Öffentlichkeitsarbeit. Als Wissenschaftler in der Abteilung für Theoretische Astrophysik der Universität Tübingen ist er im Sonderforschungsbereich/Transregio »Gravitationswellenastronomie« tätig. Seine Schwerpunkte sind Schwingungsmoden nicht-rotierender schwarzer Löcher und Neutronensterne sowie Grundlagen und Anwendungen wissenschaftlicher Visualisierung in Astrophysik und Relativitätstheorie. Für den Bereich Öffentlichkeitsarbeit des SFB konzipiert und betreut er die Ausstellung »Einstein-Wellen-Mobil« sowie diverse öffentliche Veranstaltungen.

Hanns Ruder, zunächst Professor für Theoretische Physik in Erlangen-Nürnberg und seit 1982 Professor für Theoretische Astrophysik an der Universität Tübingen, jetzt im Ruhestand, beschäftigte sich schwerpunktmäßig mit Materie in extrem starken Magnetfeldern, Radiopulsaren, Röntgenpulsaren, magnetisierten Weißen Zwergen und Neutronensternen, numerischen Methoden in der Hydrodynamik und

Allgemeinen Relativitätstheorie, Akkretionsscheiben und -säulen, Biomechanik und Visualisierung. Dafür wurde er mit zahlreichen Preisen ausgezeichnet, darunter 2006 mit der Medaille für Naturwissenschaftliche Publizistik der Deutschen Physikalischen Gesellschaft. Seine Arbeitsgruppe konzipierte das Einstein-Mobil, eine Wanderausstellung, die Schülerinnen und Schülern in ganz Deutschland die Relativitätstheorie näher bringt.

Nina Schaller, geboren 1974 in Kronberg (Taunus), studierte Biologie in Heidelberg. Zur Vorbereitung auf ihre Doktorarbeit mit handaufgezogenen Straußen absolvierte sie ein Volontariat in der Tierpflege im Zoo Frankfurt am Main und wurde Mitarbeiterin in der Abteilung BioloKomotion am Frankfurter Forschungsinstitut Senckenberg. Für ihre interdisziplinäre Doktorarbeit musste sie methodisches Neuland betreten und ein Netzwerk von Kooperationen aufbauen, u.a. mit den Universitäten in Heidelberg und Antwerpen. 2008 wurde sie in Heidelberg promoviert. Seit 2009 ist sie Postdoktorandin am *Royal Ontario Museum* und der Universität Toronto, Kanada.

Ausgezeichnet wurde sie 2009 mit dem Klaus Tschira Preis für verständliche Wissenschaft. Gefördert von der Klaus Tschira Stiftung, rief sie außerdem eine Sommerakademie für naturwissenschaftlich interessierte Teenager aus dem Großraum Rhein-Neckar ins Leben.

A. Julia Stähler, geboren 1978 in Berlin, studierte Physik an der Freien Universität Berlin, wo sie 2007 auch promovierte. Bei mehreren Forschungsaufenthalten im Ausland vertiefte und erweiterte sie ihre methodischen Kenntnisse im Bereich zeitaufgelöster Spektroskopie, beispielsweise 2008/09 als Postdoktorandin an der Universität Oxford. Am Fritz-Haber-Institut der Max-Planck-Gesellschaft leitet sie seit 2009 eine Arbeitsgruppe zur Elektronendynamik an Grenzflächen und in Festkörpern, die mit femtosekunden-zeitaufgelöster nicht-linear-optischer Spektroskopie und Zweiphotonen-Photoemission untersucht wird. Seit 2011 ist sie Leiterin eines Teilprojekts des Sonderforschungsbereichs 951 der Deutschen Forschungsgemeinschaft.

A. Julia Stähler erhielt 2008 den Klaus Tschira Preis für verständliche Wissenschaft sowie das Feodor-Lynen-Forschungsstipendium der Alexander-von-Humboldt-Stiftung, welche sie 2011 auch im Rahmen ihres Alumni-Programms als Gastwissenschaftlerin an der *Columbia University* (New York) förderte.

Martin Uecker, promovierter Physiker, ist Mitarbeiter von Professor Frahm in der Biomedizinischen NMR Forschungs GmbH am Max-Planck-Institut für biophysikalische Chemie in Göttingen. Schwerpunkt seiner Forschungen ist die Weiterentwicklung von Algorithmen für die Bild-Rekonstruktion in der Magnetresonanz-Tomografie. Dies gilt insbesondere für Verfahren der parallelen Datenaufnahme, der nicht-kartesischen MRT und der Echtzeit-MRT.

Dirk Voit, promovierter Physiker, ist Mitarbeiter von Professor Frahm in der Biomedizinischen NMR Forschungs GmbH am Max-Planck-Institut für biophysikalische Chemie in Göttingen. Seine Arbeiten befassen sich mit der hochauflösenden Magnetresonanz-Tomografie von Hirnfunktionen, der radialen Echtzeit-MRT, der quantitativen Flussmessung und der Trennung von Wasser- und Fettanteilen in den MRT-Bildern.

Shuo Zhang, promovierter Physiker, ist Mitarbeiter von Professor Frahm in der Biomedizinischen NMR Forschungs GmbH am Max-Planck-Institut für biophysikalische Chemie in Göttingen. Seine aktuellen Forschungen konzentrieren sich auf die Weiterentwicklung der Echtzeit-MRT für Gelenkuntersuchungen, Darstellungen der Schluck- und Sprechvorgänge, sowie kardiovaskuläre Anwendungen.

Redaktion: Almut Popp

Bibliografische Information der Deutschen Nationalbibliothek

Die Deutsche Nationalbibliothek verzeichnet
diese Publikation in der Deutschen Nationalbibliografie;
detaillierte bibliografische Daten sind im Internet
über <http://dnb.d-nb.de> abrufbar.

© Wallstein Verlag, Göttingen 2011
www.wallstein-verlag.de

Vom Verlag gesetzt aus der Stempel Garamond
Umschlaggestaltung: Basta Werbeagentur, Steffi Riemann
Druck: Friedrich Pustet KG, Regensburg
ISBN 978-3-8353-1027-8