

FORSCHUNGS FORUM

PADERBORN

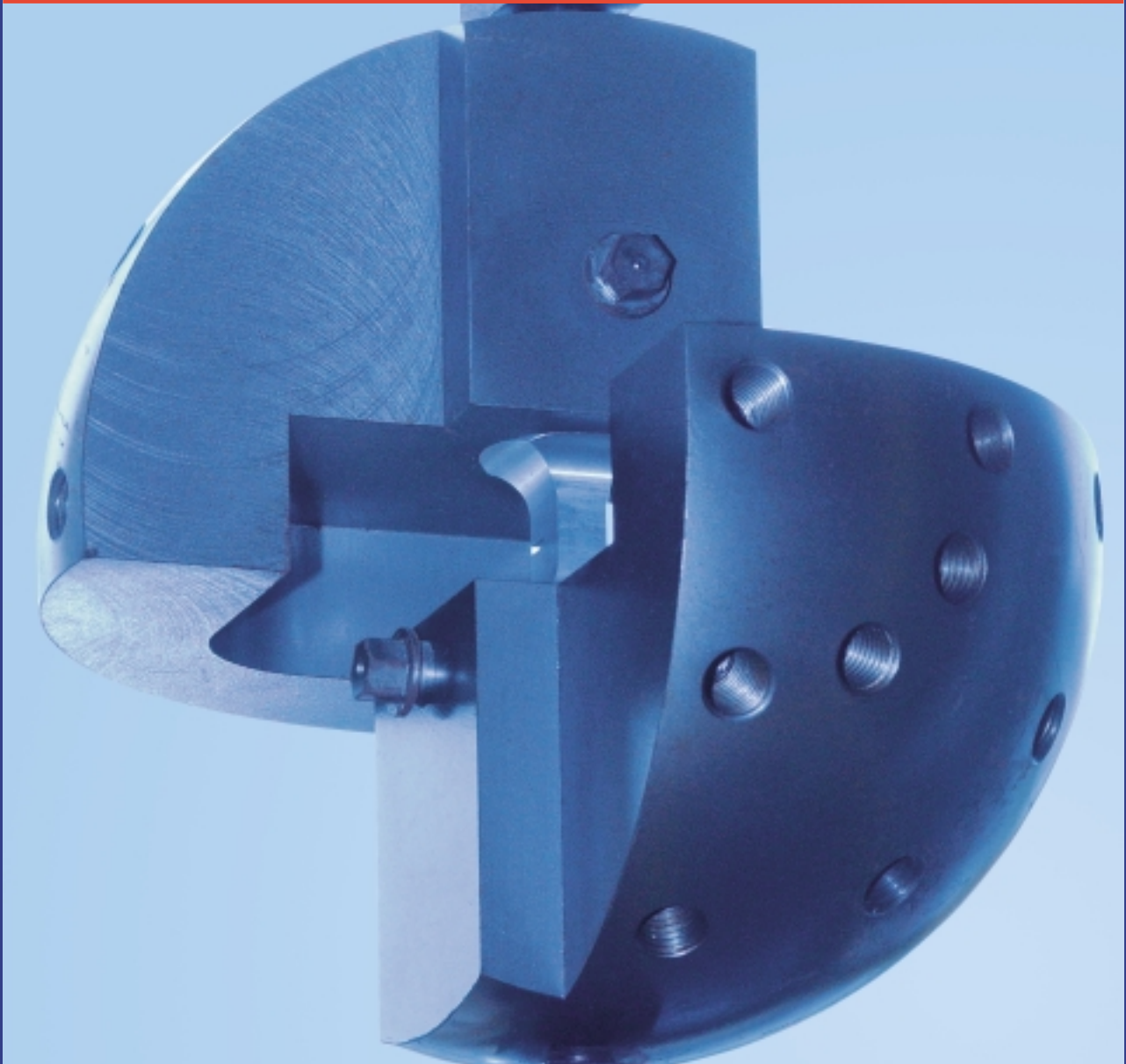


UNIVERSITÄT PADERBORN

Die Universität der Informationsgesellschaft

8-2005

P A D E R B O R N E R U N I V E R S I T Ä T S M A G A Z I N



- | | |
|---|---|
|  Das W-LAN der Zukunft |  High-Tech-Unternehmensgründungen |
|  Neue Gesetze des Lichts |  ICE-Radreifenbruch |
|  Rationalismus zwischen Vision und Wirklichkeit |  Schnelle Mikrochips |

Feiern Sie Erfolge in **park**-Lage!

Richtiger
Mehrwert für Sie!

Plus-Pluspunkte im Park

Erstklassiges Image
Innovatives Umfeld
Optimale
Verkehrsanbindung
Vielfältige Services
Gelebter Know how-Transfer

- **Tagungsräume**
(für 10-200 Personen)
- **Mietflächen**
(für Büro und Labor)
- **Grundstücke**
(für Neubauten)

TechnologiePark Paderborn GmbH

Technologiepark 13 · 33100 Paderborn

Fon 0 52 51 / 1 60 90-10

Fax 0 52 51 / 1 60 90-49

Mail: info@technologiepark-paderborn.de

www.technologiepark-paderborn.de

Kooperationspartner der Universität Paderborn



IMPRESSUM

Herausgeber

Prof. Dr. Nikolaus Risch
Rektor der Universität Paderborn

Konzeption und Redaktion

Ramona Wiesner
Leiterin des Referats Hochschulmarketing
und Universitätszeitschrift
Warburger Str. 100, 33098 Paderborn
Tel.: 05251/60 2553, 3880
E-Mail: wiesner@zv.uni-paderborn.de
<http://www.zit.uni-paderborn.de/hochschulmarketing/>

ForschungsForum Paderborn (ffp) im Internet

<http://www.uni-paderborn.de/ffp/>

Wissenschaftlicher Beirat

Prof. Dr. Gitta Domik
Prof. Dr. Jörg Jarnut
Prof. Dr. Klaus Meerkötter
Prof. Dr. Winfried Reiß
Prof. Dr. Wilhelm Schäfer
Prof. Dr. Jürgen Voß
Prof. Dr. Jörg Wallaschek
Prof. Dr. Gerhard Wortmann

Drucklegung

Januar 2005

ISSN (Print) 1435-3709

Layout

PADA-Werbeagentur
Heierswall 2, 33098 Paderborn

Anzeigenverwaltung

PADA-Marketingverlag
Heierswall 2, 33098 Paderborn
Tel.: 05251/527577

Auflage

5 000

Titel

Patentierte Belastungsvorrichtung für mehrachsige Bruchversuche im Bereich der Angewandten Mechanik. Siehe auch Bericht zum ICE-Radreifenbruch.



Editorial



Ramona Wiesner
Referentin für Öffentlichkeitsarbeit

Liebe Leserinnen und Leser,

ist es Magie oder Wirklichkeit? Bremse einen Lichtstrahl ab, knicke, biege, teile, manipulierte ihn oder nimm eine hell strahlende Lichtquelle und baue einen durchsichtigen Käfig um sie herum, der keinen einzigen Strahl entweichen lässt! Exakt dies haben sich Physiker aus aller Welt vorgenommen: Die bisher bekannten Gesetze des Lichts umzuschreiben. Wissenschaftler wollen Licht in Zukunft beeinflussen können (Seite 10).

Der technologische Fortschritt wird nicht nur durch neue Hardware vorangetrieben, sondern auch durch bessere Software. Ein Beispiel sind mobile Ad-hoc-Netzwerke, drahtlose Kommunikationsnetzwerke, die im Sonderforschungsbereich 376 Massive Parallelität, Teilprojekt C6 Mobile Ad-hoc-Netzwerke am Institut für Informatik an der Universität Paderborn erforscht werden (Seite 38).

Unfallursache für den 1998 verunglückten ICE in Eschede war der gebrochene Radreifen eines gummigefederten Rades. Die Radumdrehungen verursachten eine Beanspruchung, die zu einem ausgedehnten Ermüdungsrisswachstum und letztlich zum Bruch des Radreifens führte. Paderborner Untersuchungen belegen, dass mit der Bruchmechanik das Risswachstum im ICE-Radreifen erklärt werden kann (Seite 44).

Die 8. Ausgabe des „ForschungsForum Paderborn“ zeigt 10 weitere Ausschnitte aus der Forschungslandschaft der Universität Paderborn. Finanziert wird das Wissenschaftsmagazin teilweise durch Werbung. Dadurch ist es möglich, Forschungsthemen der Universität Paderborn in anspruchsvoller Form vorzustellen. Wir danken allen herzlich, die unser Projekt unterstützt haben.

Ihnen, liebe Leserinnen und Leser, wünsche ich wie immer viele neue Erkenntnisse beim Lesen des Forschungsmagazins.

Ihre Ramona Wiesner

Seite 6

Wohin geht die Reise?

Einstellungen der Bundesbürger zu Urlaubsreisen

Prof. Dr. phil. Albrecht Steinecke

Fakultät für Kulturwissenschaften



Seite 10

Neue Gesetze des Lichts

Die Idee der photonischen Kristalle

Prof. Dr. rer. nat. Dipl.-Phys. Ralf B. Wehrspohn, apl. Prof. Dr. rer. nat. Dipl.-Phys. Siegmund Greulich-Weber

Fakultät für Naturwissenschaften



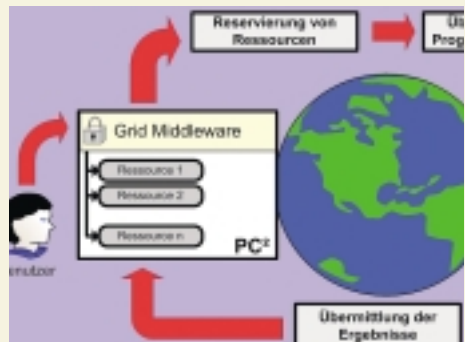
Seite 16

Grid Computing

Kooperation und gemeinsame Nutzung vernetzter Ressourcen

Prof. Dr. rer. nat. Odej Kao, Dipl.-Inform. Matthias Hovestadt

Fakultät für Elektrotechnik, Informatik und Mathematik



Seite 22

Vermögenssteuer, Verluste und Investitionsentscheidungen

Ökonomische Analyse der Renaissance einer umstrittenen Steuer

Prof. Dr. rer. pol. Dipl.-Kffr. Caren Sureth

Fakultät für Wirtschaftswissenschaften



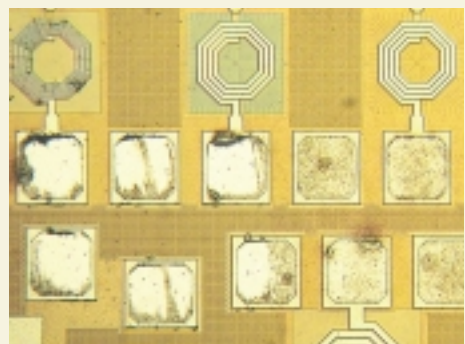
Seite 28

Schnelle Mikrochips

Preiswerter Anschluss an die Datenautobahn

Prof. Dr.-Ing. Andreas Thiede, M.Sc. Zheng Gu

Fakultät für Elektrotechnik, Informatik und Mathematik



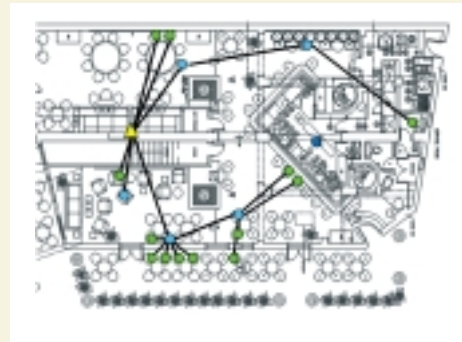
Seite 34

Rationalismus zwischen Vision und Wirklichkeit
Descartes, Leibniz und die Grenzen des Machbaren
 Prof. Dr. phil. Volker Peckhaus
 Fakultät für Kulturwissenschaften



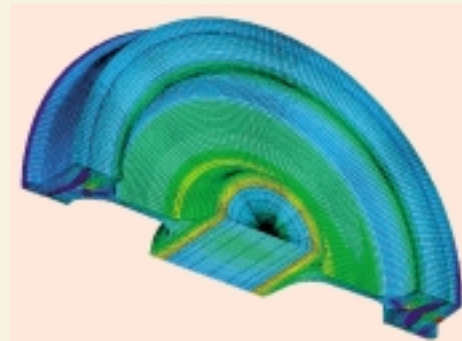
Seite 38

Mobile Ad-hoc-Netzwerke
Das W-LAN der Zukunft
 PD Dr. rer. nat. Christian Schindelhauer
 Fakultät für Elektrotechnik, Informatik und Mathematik



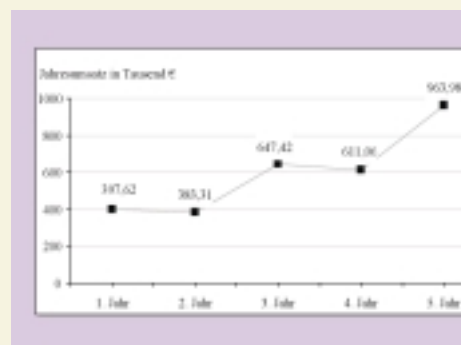
Seite 44

ICE-Radreifenbruch
Bruchmechanik trägt zur Klärung der Schadensursache bei
 Prof. Dr.-Ing. Hans Albert Richard
 Fakultät für Maschinenbau



Seite 50

High-Tech-Unternehmensgründungen
Wer überlebt und wer überlebt nicht?
 Prof. Dr. rer. pol. habil, Dr. h.c. mult., Dipl.-Kfm. Wolfgang Weber,
 Dipl.-Kffr. Anja Schmelter
 Fakultät für Wirtschaftswissenschaften



Seite 56

Viel Licht und keine Blendung
Aktive Scheinwerfersysteme zur Erhöhung der Verkehrssicherheit
 Prof. Dr.-Ing. Jörg Wallaschek, Dipl.-Ing. Rainer Kauschke,
 Dipl.-Ing. Jacek Roslak
 Fakultät für Maschinenbau





Das **«** Wir machen den Weg frei **»** Prinzip

200 Tausend Kunden, 90 Tausend Mitglieder, ein Prinzip.

Das Ergebnis: *Einzigartige Kundennähe*



**Volksbank
Paderborn-Höxter**
mit uns zum Erfolg

www.vb-paderborn-hoexter.de



Prof. Dr. Wilhelm Schäfer

Vorwort

Liebe Leserin, lieber Leser,

Exzellenz in der Forschung ist ein zurzeit in der Öffentlichkeit vor allem unter dem Stichwort „Eliteuniversität“ breit diskutiertes Thema. Die Bundesregierung plant gemeinsam mit den Bundesländern ein 1,9 Mrd. Euro Programm zum Ausbau von Spitzenforschung in den Jahren 2006 bis 2011. Allerdings sind die Fallstricke zur Einrichtung dieses Programms noch lange nicht beseitigt. Die Gelder sollen nicht aus anderen Budgets, insbesondere im Bereich der Bildung, abgezogen werden, sondern durch Streichung der Eigenheimzulage gewonnen werden. Außerdem entstehen durch die Beteiligung der Länder Diskussionen, die im Rahmen der Föderalismusdebatte mit anderen Streitpunkten zwischen Bundes- und Landesregierungen verquickt werden.

Einhellig begrüßt die Wissenschaft in Deutschland ein solches Programm, um die langfristige wirtschaftliche und wissenschaftliche Stärke des Standorts „Deutschland“ zu behaupten. Vergleicht man Ausgaben für „Wissenschaft und Bildung“ in allen Ländern der Erde, wird immer wieder deutlich, dass Deutschland in diesem Bereich immer noch „hinterherhinkt“.

Problematisch sehen viele Wissenschaftsorganisationen, aber auch einzelne Wissenschaftler, dass dieses Programm zu sehr unter dem Thema „Eliteuniversität“ diskutiert wird und dass sogar schon erste Listen mit 10 Vorschlägen für die entsprechenden Standorte existieren. Hier stellt sich die Frage, ob „Elite“ in dieser Weise gemessen werden kann, ob Elite staatlich verordnet werden kann und ob es wirklich Universitäten gibt, die in allen Fächern „Elite“ sind. Selbst die immer wieder im Vergleich angeführten US-amerikanischen Standorte, wie Harvard, Yale oder Stanford, sind durch Spitzenleistung nur in einzelnen Fächern bekannt geworden, haben ihren Status über Jahrzehnte erworben und verfügen darüber hinaus noch über bei weiterem höhere Finanzmittel als selbst das oben angesprochene Programm der Bundesregierung den „deutschen Eliteuniversitäten“ einbringen würde. (Harvard allein verfügt beispielsweise über ein Stiftungsvermögen von ca. 26 Mrd. USD).

Wesentlich sinnvoller erscheinen die beiden anderen ebenfalls in dem Programm geplanten Forschungsvorhaben zu sein, die nicht die Einrichtung dieser so genannten „deutschen Eliteuniversitäten“ zum Ziel haben. Mit einem Finanzvolumen von ca. 6–8 Mill. Euro pro Jahr bzw. 1 Mill. Euro pro Jahr sollen an einzelnen Hochschulen so genannte „Exzellenzcluster“ bzw. Graduiertenschulen, das sind Promotionsprogramme für international ausgesuchte besonders qualifizierte Studierende, eingerichtet werden.

„Exzellenzcluster“ sollen einen oder mehrere bereits ausgewiesene Forschungsschwerpunkte einer Universität weiter verstärken und ausbauen. Ausbau soll, soweit die Antragsrichtlinien bisher bekannt sind, insbesondere auch die Einbindung außeruniversitärer Forschung in Kooperation mit (regionaler) Industrie zum Schwerpunkt haben.

Die Universität Paderborn wird sich im Bereich ihres Schwerpunkts Informatik und ihre Anwendungen in den Ingenieurwissenschaften an dieser Ausschreibung beteiligen. Hier ist sie durch zwei Sonderforschungsbereiche, sowie Graduiertenkollegs der Deutschen Forschungsgemeinschaft, die International Graduate School des Landes NRW, das Heinz Nixdorf Institut, sowie eine Reihe von Forschungsk Kooperationen, wie das L-LAB (mit der Fa. Hella), das C-LAB (mit der Fa. Siemens), sowie zwei Einrichtungen der Fraunhofer-Gesellschaft hervorragend positioniert. Chancen hätten vielleicht auch Anträge zu Graduiertenschulen aus den Bereichen Kultur-, Natur- und Wirtschaftswissenschaften.

Kristallisationspunkte aller dieser Bemühungen sind die Ergebnisse unserer Forscherinnen und Forscher, von denen Sie einige in dieser Ausgabe des ForschungsForums finden. Ich wünsche Ihnen viel Spaß und vielleicht die eine oder andere Anregung beim Lesen.

Wilhelm Schäfer

Prorektor für Forschung und wissenschaftlichen Nachwuchs

Wohin geht die Reise?

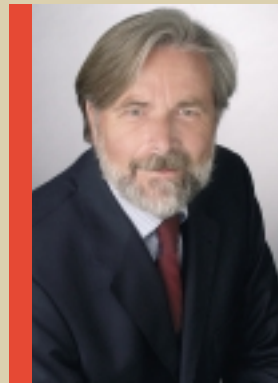
Einstellungen der Bundesbürger zu Urlaubsreisen

Seit den 1990er Jahren vollzieht sich innerhalb des bundesdeutschen und internationalen Tourismusmarktes ein gravierender und rapider Strukturwandel. Den Hintergrund bildet dabei die angespannte Wettbewerbssituation, die durch mehrere Faktoren ausgelöst wird: Dazu zählen u. a. die wachsenden Ansprüche der Konsumenten, das Auftreten neuer regionaler Wettbewerber, die Konzentrationstendenzen bei Reiseveranstaltern sowie die innovativen Vertriebsformen (vgl. STEINECKE 2000).

Dieser Strukturwandel wurde durch die Terroranschläge des 11. Septembers 2001 in New York noch verstärkt: Seitdem ist nämlich – zum ersten Mal seit 1982 – ein genereller Rückgang des Nachfrageaufkommens im internationalen Tourismus zu beobachten, der regional allerdings sehr unterschiedlich verläuft. Besonders die USA, aber auch einige islamische Staaten sind deutlich stärker betroffen als Zielgebiete in Europa und Asien (vgl. WTO 2002). Gegenüber anderen krisenhaften Erscheinungen (z. B. Golfkrieg, Anschläge in Sri Lanka und Ägypten), die den Tourismusmarkt nur kurzzeitig bzw. nur auf nationaler Ebene beeinflusst haben, scheinen die Nachwirkungen des 11. Septembers 2001 von längerer Dauer zu sein und einen globalen Charakter zu haben. Gründe dafür liegen u. a. in weiteren Terroranschlägen, bei denen Touristen die Opfer waren (z. B. Insel Djerba, Bali, Kenia), aber auch im Irak-Krieg und seinen anhaltenden Nachwirkungen.

Fragestellung und Methodik der Untersuchung

Vor dem Hintergrund der skizzierten Dynamik des Tourismusmarktes sollten im Rahmen einer empirischen Untersuchung – im Sinne einer Zwischenbilanz – die aktuellen Einstellungen der deutschen Bevölkerung zu Urlaubsreisen erfasst werden. Im



Prof. Dr. phil. Albrecht Steinecke ist seit dem Wintersemester 1997/98 Inhaber des Lehrstuhls für Wirtschafts- und Fremdenverkehrsgeographie. Seine Arbeitsschwerpunkte sind touristische Trendforschung, Kulturtourismus, Erlebnis- und Konsumwelten sowie touristische Zielgruppen- und Destinationsanalysen.

Mittelpunkt stand dabei die Frage, ob die Bundesbürger nach einer Phase der Verunsicherung ihr bisheriges Reiseverhalten voraussichtlich wieder aufnehmen oder ob sich die erkennbaren Veränderungen im bundesdeutschen Tourismus stabilisieren werden. Innerhalb der Tourismusbranche werden derzeit vor allem folgende Trends diskutiert: Ein angeblich sinkendes Interesse an Urlaubsreisen, deutliche Buchungsrückgänge in traditionellen Warmwasser-Destinationen, ein Boom von Niedrigpreis-Destinationen in Osteuropa, die Bedeutungszunahme von erdgebundenen Zielen, ein wachsendes Interesse an Last-Minute-Reisen, die generelle Schnäppchenmentalität, der Boom der No-Frills-Carrier, ein kurzfristiges Buchungsverhalten etc. (vgl. F.U.R. 2003, S. I).

Die Studie wurde durch die Forschungsgruppe „Tourismus“ an der Universität Paderborn erstellt, der Prof. Dr. Albrecht Steinecke (Paderborn), Dr. Wolfgang Isenberg (Bensberg) und Dr. Horst Martin Mullenmeister (Hannover) angehören. Als Sponsor der Untersuchung fungierte die TUI AG (Hannover). Inhaltliche Schwerpunkte der empirischen Erhebung waren Einstellungen der Bundesbürger zu Urlaubsreisen generell, Erwartungen an Urlaubsreiseziele, Verhaltensweisen während der letzten Urlaubsreise sowie Einstellungen zu künftigen Urlaubsreisen (vgl. ISENBERG/MÜLLENMEISTER/STEINECKE 2003). Zur Beantwortung dieser Forschungsfragen war es notwendig, über eine aussagekräftige und aktuelle Datenbasis zu verfügen. Aus diesem Grund hat die Forschungsgruppe „Tourismus“ im November 2002 die Firma TNS EMNID (Bielefeld) beauftragt, eine repräsentative demoskopische Untersuchung durchzuführen (1 998 computergestützte Telefoninterviews in zwei Wellen). Als Grundgesamtheit fungierte dabei die Wohnbevölkerung in Privathaushalten im Alter von 14 und mehr Jahren in der Bundesrepublik Deutschland.

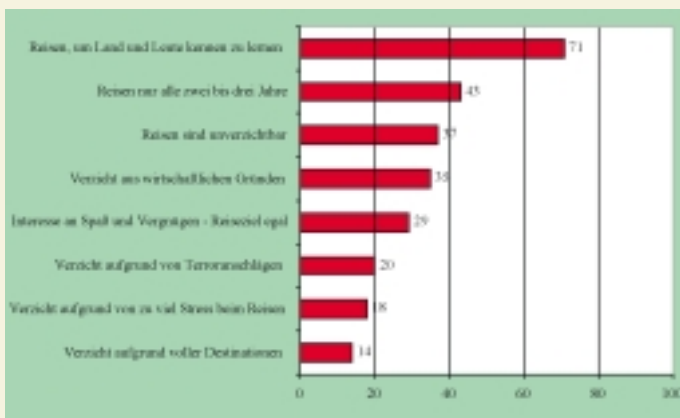


Abb. 1: Einstellungen der Deutschen zu Urlaubsreisen generell (in Prozent). n = 1 998; Zustimmung zu Statements (trifft voll und ganz zu, trifft eher zu)



Foto: A. Steinecke

Das Interesse an Land und Leuten ist bei den bundesdeutschen Touristen deutlich ausgeprägter als die ausschließliche Spaßorientierung (unabhängig vom Reiseziel). Speziell historische Städte, Kulturdenkmäler und kulturelle Events (z. B. die Bregenzer Festspiele) verzeichnen eine wachsende Nachfrage.

Einstellungen zu Urlaubsreisen generell

Obwohl die Sicherheitsproblematik auf Reisen derzeit in der öffentlichen Diskussion im Vordergrund steht, ist generell davon auszugehen, dass die Einstellungen zu Urlaubsreisen (wie auch das Reiseverhalten) nicht ausschließlich von Sicherheitsüberlegungen, sondern vom Zusammenwirken mehrerer Steuerfaktoren bestimmt werden: Neben soziodemographischen Merkmalen zählen dazu die persönliche wirtschaftliche Lage, bisherige Reiseerfahrungen, grundsätzliche Reisemotive etc. (vgl. KULNAT 2003). Aus diesem Grund wurden die generellen Einstellungen der Bundesbürger zu Urlaubsreisen mit Hilfe von mehreren Statements erfasst:

- „Auf eine Urlaubsreise im Jahr kann ich auf keinen Fall verzichten, ich spare lieber an anderen Dingen.“
- „Mir reicht es aus, alle zwei bis drei Jahre eine Urlaubsreise zu unternehmen.“
- „In absehbarer Zeit kann ich mir aus wirtschaftlichen Gründen keine längere Urlaubsreise leisten, allenfalls Tagesausflüge oder Wochenendreisen.“
- „Aufgrund der Terroranschläge in den letzten Monaten werde ich im kommenden Jahr auf eine Urlaubsreise verzichten.“
- „Urlabsreisen sind für mich mit zu viel Stress verbunden, ich bleibe im kommenden Jahr deshalb in meinem Urlaub lieber zu Hause.“
- „In den Urlaubsregionen ist es viel zu voll und zu laut, deshalb unternehme ich künftig keine Urlaubsreisen mehr.“
- „Bei Urlaubsreisen will ich vor allem Spaß und Vergnügen – das Reiseziel ist mir egal.“
- „Meinen Urlaub nutze ich gern, um meinen Horizont zu erweitern und um andere Länder und Völker kennen zu lernen.“

Gegenwärtig lässt sich in der bundesdeutschen Bevölkerung folgendes Einstellungsspektrum zu Urlaubsreisen feststellen (vgl. Abbildung 1):

- Interesse an Land und Leuten dominiert deutlich vor der Spaßorientierung.

An erster Stelle aller Nennungen rangiert das Interesse, im Urlaub den Horizont zu erweitern sowie andere Länder und Völker kennen zu lernen (71 Prozent), während die entgegengesetzte Haltung – eine ausschließliche Spaßorientierung, bei der das Reiseziel keine Bedeutung hat – deutlich geringer ausgeprägt ist (29 Prozent). Ein besonders großes Interesse an Land und Leuten findet sich bei Bundesbürgern mit folgenden Merkmalen: 50- bis 59-Jährige, Personen mit hohem Bildungsstand und hohem Einkommen sowie Regelmäßig-Reisende. Bei den Bundesbürgern, die ein besonders großes Interesse an Spaß und Vergnügen haben und denen das Reiseziel egal ist, handelt es sich vor allem um Jugendliche und junge Erwachsene, um Männer und um Intervall-Urlauber.

- Intervall-Reisende stellen innerhalb der bundesdeutschen Gesellschaft eine größere Gruppe dar als die Regelmäßig-Reisenden.

Für 37 Prozent der Befragten stellen Urlaubsreisen ein zentrales Konsumgut dar: Sie sparen eher an anderen Dingen, als auf eine Urlaubsreise zu verzichten. Noch größer ist allerdings der Anteil von Bundesbürgern, denen es ausreicht, alle zwei bis drei Jahre eine Urlaubsreise zu unternehmen (43 Prozent). Der Wunsch nach einer Urlaubsreise pro Jahr ist besonders stark ausgeprägt bei 30- bis 49-Jährigen, Personen mit mittlerer und höherer Schulbildung sowie hohem Einkommen, Regelmäßig-Reisenden und „Land und Leute“-Urlaubern. Zu den Merkmalen der Bundesbürger, die sich in besonderem Maß als Intervall-Reisende

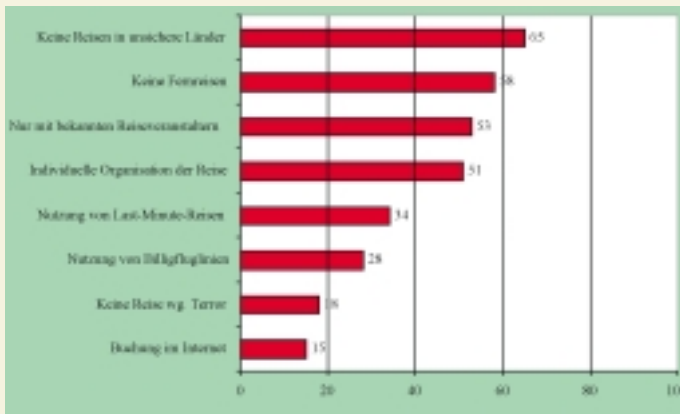


Abb. 2: Einstellungen der Deutschen zu zukünftigen Urlaubsreisen (in Prozent).
n = 1 998; Zustimmung zu Statements (trifft voll und ganz zu, trifft eher zu)

betrachten, zählen hingegen: Niedrige Schulbildung sowie unteres und mittleres Einkommen, geringe Reishäufigkeit seit 1997 und Spaßurlauber.

- Wirtschaftliche Gründe spielen eine größere Rolle bei einem eventuellen Reiseverzicht als Sicherheitsüberlegungen.

35 Prozent der Befragten haben angegeben, dass sie sich in absehbarer Zeit aus wirtschaftlichen Gründen keine längere Urlaubsreise leisten können, allenfalls Tagesausflüge oder Wochenendreisen. Aufgrund der Terroranschläge wollen hingegen nur 20 Prozent der Befragten künftig auf eine längere Urlaubsreise verzichten. Der Reiseverzicht aus Sicherheitsüberlegungen findet sich besonders bei Frauen, Senioren (60+ Jahre), Personen mit niedriger Schulbildung und niedrigem Einkommen sowie einer geringen Reishäufigkeit seit 1997. Aus wirtschaftlichen Gründen werden vor allem Befragte mit diesen Merkmalen in absehbarer Zeit keine Urlaubsreise unternehmen: Frauen, 14- bis 29- und 40- bis 49-Jährige, Personen mit niedriger Schulbildung und geringem Einkommen sowie geringer Reishäufigkeit seit 1997 und Spaßurlauber. Die Ergebnisse zeigen, dass diese beiden Gruppen über ähnliche Merkmale wie die Gruppe der Intervall-Reisenden verfügen: Es handelt sich jeweils um Unterschichtangehörige (Bildung, Einkommen), die sich aufgrund ihrer wirtschaftlichen Lage nur gelegentlich eine Urlaubsreise leisten können. Die gegenwärtige Sicherheitsdiskussion bietet ihnen offenbar ein zusätzliches Argument dafür, den Verzicht auf eine Urlaubsreise zu begründen.

- Nur ein geringer Teil der Bevölkerung will aufgrund von Stress auf Reisen bzw. von überfüllten Destinationen auf eine Urlaubsreise verzichten.

Die Zustimmungswerte zu den beiden Statements „Urlaubsreisen sind für mich mit zu viel Stress verbunden, ich bleibe im kommenden Jahr in meinem Urlaub lieber zu Hause“ (18 Prozent) und „In den Urlaubsregionen ist es viel zu voll und zu laut, deshalb unternehme ich künftig keine Urlaubsreisen mehr“ (14 Prozent) sind sehr niedrig. Von einer generellen Reisemüdigkeit der bundesdeutschen Bevölkerung bzw. einer negativen Haltung zu Urlaubsreisen oder Urlaubszielen kann also nicht die Rede sein. Die Befragten, die diesen beiden Statements in besonders ausgeprägtem Maß zugestimmt haben, weisen wiederum Merkmale auf, die bereits bei den Intervall-Reisenden zu finden waren: Senioren, Personen mit niedriger Schulbildung und geringem Einkommen sowie Personen mit geringer Reishäufigkeit seit 1997.

Einstellungen zu künftigen Urlaubsreisen

Neben den generellen Einstellungen zu Urlaubsreisen wurden im Rahmen der Untersuchung auch die Einstellungen der Bundesbürger zu künftigen Urlaubsreisen ermittelt (unabhängig davon, ob die Befragten für das kommende Jahr eine Urlaubsreise planen oder nicht). Die Erfassung erfolgte mit Hilfe von acht Statements:

- „Ich werde künftig nur mit bekannten Reiseveranstaltern und Fluggesellschaften verreisen.“
- „Ich werde in absehbarer Zeit nicht in bestimmte Länder reisen, die mir unsicher erscheinen.“
- „Ich werde meine Urlaubsreise in absehbarer Zeit selbst organisieren, also ohne die Hilfe von Reisebüros oder Reiseveranstaltern.“
- „Ich werde in absehbarer Zeit keine Fernreise – z. B. nach Afrika, Asien, Australien oder in die USA – unternehmen.“
- „Ich werde künftig vor allem die Angebote der Billigfluglinien nutzen.“
- „Ich werde in absehbarer Zeit aus Furcht vor Terroranschlägen keine Urlaubsreisen unternehmen.“
- „Ich werde künftig meine Urlaubsreisen überwiegend im Internet buchen.“
- „Ich werde kurzfristig entscheiden, ob ich im kommenden Jahr eine Urlaubsreise unternehmen werde und dann ein Last-Minute-Angebot nutzen.“

Hinsichtlich der künftigen Urlaubsreisen lassen sich bei den Bundesbürgern folgende Einstellungen beobachten (vgl. Abbildung 2):

- Nur ein geringer Teil der Bundesbürger will aufgrund der Terroranschläge in absehbarer Zeit auf Urlaubsreisen verzichten, doch viele Deutsche wollen keine Reisen in unsichere Länder bzw. keine Fernreisen unternehmen und auch nur mit bekannten Reiseveranstaltern bzw. Fluggesellschaften verreisen.

Nur jeder fünfte Befragte gab an, dass er aufgrund der Terroranschläge in absehbarer Zeit auf eine Urlaubsreise verzichten wird. Jeweils eine hohe Zustimmung erfahren allerdings die vier Statements, die sich auf Sicherheitsaspekte bei Urlaubsreisen beziehen: Offenbar werden die Bundesbürger nicht auf Urlaubsreisen verzichten, aber ihr künftiges Reiseverhalten ändern. Als unsicher betrachtete Länder werden künftig vor allem gemieden von Frauen, von 30- bis 59-Jährigen, von mittleren Bildungs- und Einkommensgruppen sowie von Regelmäßig-Reisenden. Die Einstellung, in absehbarer Zeit keine Fernreise zu unternehmen, findet sich besonders bei Frauen, bei 40- bis 59-Jährigen, bei mittleren Bildungs- und höheren Einkommensgruppen sowie bei Intervall-Reisenden. Mit bekannten Reiseveranstaltern und Fluggesellschaften verreisen wollen künftig vor allem Ostdeutsche, 30- bis 39-Jährige, Personen mit mittlerer Bildung, Regelmäßig-Reisende sowie „Land und Leute“-Urlauber. Aufgrund der Terroranschläge wollen besonders folgende Gruppen in absehbarer Zeit auf eine Urlaubsreise verzichten: Frauen, Senioren, Personen mit niedriger Bildung und geringem Einkommen, Befragte mit einer geringen Reishäufigkeit seit 1997 und generell negativen Einstellungen zum Reisen (Reiseverzicht aus wirtschaftlichen Gründen, aus Stress bzw. wegen überfüllter Zielgebiete).



Thematisierte Erlebnis- und Konsumwelten (wie das Shopping-Center „Mercato“ in Dubai im Stil der italienischen Renaissance) erleben seit den 1990er-Jahren weltweit einen Boom. Auch bei den bundesdeutschen Touristen erfreuen sich multifunktionale Freizeiteinrichtungen einer großen Beliebtheit: Jeder Dritte hat im Urlaub vor allem Interesse an Spaß und Vergnügen, während ihm das Reiseziel egal ist.

- **Großes Interesse besteht an einer individuellen Organisation der künftigen Urlaubsreise, während neue Angebots- und Vertriebsformen (Last-Minute-Reisen, Billigfluglinien und Buchungen im Internet) bislang erst für einen kleinen Teil der bundesdeutschen Bevölkerung von Interesse sind.**

Besonders interessiert an einer individuellen Organisation der künftigen Urlaubsreise waren 40- bis 59-Jährige, Personen mit einer höheren Bildung und einem höheren Einkommen sowie Regelmäßig-Reisende. Die Zustimmung zu den Statements, die sich auf neue Angebots- und Vertriebsformen im Tourismus beziehen, ist – auf alle Befragten und somit auf die gesamte deutsche Bevölkerung bezogen – relativ niedrig. Ein besonders ausgeprägtes Interesse an der Nutzung von Billigfluglinien haben Männer, Jugendliche und junge Erwachsene sowie Spaßurlauber. Ähnliche Merkmale weist die Gruppe der Befragten auf, die ihre Reisen künftig überwiegend im Internet buchen wollen: Es handelt sich um Männer, um Jugendliche und junge Erwachsene, um höhere Bildungs- und Einkommensgruppen sowie um Regelmäßig-Reisende. Eine besonders kurzfristige Reiseentscheidung und ein großes Interesse an Last-Minute-Reisen finden sich bei Ostdeutschen, bei Jugendlichen und jungen Erwachsenen, bei mittleren Bildungsgruppen, bei Personen mit mittlerem und höherem Einkommen, bei Intervall-Reisenden und bei spaßorientierten Befragten.

Fazit:

Tourismus im Wandel

- **Die meisten Urlauber wollen Inhalte – nicht nur Strand und Sonne.**

Die Analyse der generellen Einstellungen zu Urlaubsreisen zeigt deutlich, dass das Interesse an „Land und Leuten“ deutlich stärker ausgeprägt ist als eine ausschließliche Spaßorientierung, unabhängig vom Reiseziel. Ein durchgängiges Interesse der Reisenden an ihrer Urlaubsumwelt lässt sich auch an anderen Ergebnissen der Studie belegen, die hier aus Platzgründen nicht dargestellt werden können: Es ist nämlich nicht nur bei der Einstellung zu Urlaubsreisen generell vorhanden, sondern auch

bei den Erwartungen an das Reiseziel festzustellen (Möglichkeit zu Kontakt mit „Land und Leuten“) und schließlich beim Verhalten während der Reise zu beobachten (eigene Aktivitäten mit Interesse am Reiseziel).

- **Sicherheit ist für Urlauber wichtig – aber nicht alles.**

In als unsicher wahrgenommenen Zeiten gehen die Bundesbürger das Thema „Sicherheit“ auf ihre Weise an. Aus Sicherheitsgründen will nur ein geringer Teil der Bevölkerung auf seine Urlaubsreise verzichten (dabei handelt es sich vorrangig um Bundesbürger mit geringer Reiseerfahrung). Hingegen spielen wirtschaftliche Gründe bei einem möglichen Reiseverzicht eine weitaus gewichtigere Rolle. Allerdings werden die Bundesbürger ihr Reiseverhalten künftig aufgrund der Sicherheitslage ändern: Sie werden keine Reisen in Länder unternehmen, die als unsicher betrachtet werden; auch Fernreisezielen stehen sie skeptisch gegenüber. Das wachsende Sicherheitsbedürfnis wird auch dazu führen, dass bevorzugt Angebote bekannter Reiseveranstalter und Fluggesellschaften genutzt werden.

- **Die generelle Reiselust der Deutschen ist ungebrochen – allerdings weisen wichtige Steuerfaktoren des Reisens gegenwärtig eine negative Entwicklung auf (Wohlstand, Zeitbudget).**

Das Interesse der Bevölkerung in Deutschland an Urlaubsreisen ist ungebrochen. Von einer generellen Reismüdigkeit bzw. einem sinkenden Interesse der Deutschen an Reisen kann gegenwärtig nicht die Rede sein. So will nur ein kleiner Teil der Befragten künftig auf Urlaubsreisen verzichten, weil – ihrer Meinung nach – Reisen mit zu viel Stress verbunden sind bzw. weil die Urlaubsregionen als zu voll und zu laut betrachtet werden. Der wichtigste Grund für einen möglichen Reiseverzicht liegt hingegen in der persönlichen wirtschaftlichen Lage. Die künftige Dynamik der touristischen Nachfrage wird somit wesentlich von der Entwicklung der bundesdeutschen Wirtschaft generell und speziell des Arbeitsmarktes abhängen. Darüber hinaus wird bei künftigen Untersuchungen auch das Problem der längeren Arbeitszeiten zu berücksichtigen sein, die gegenwärtig wieder eingeführt werden: In der Vergangenheit zählten nämlich die sinkenden Arbeitszeiten – neben dem breiten Wohlstand und der wachsenden Motorisierung – zu den wichtigen Motoren des anhaltenden touristischen Booms.

Literatur

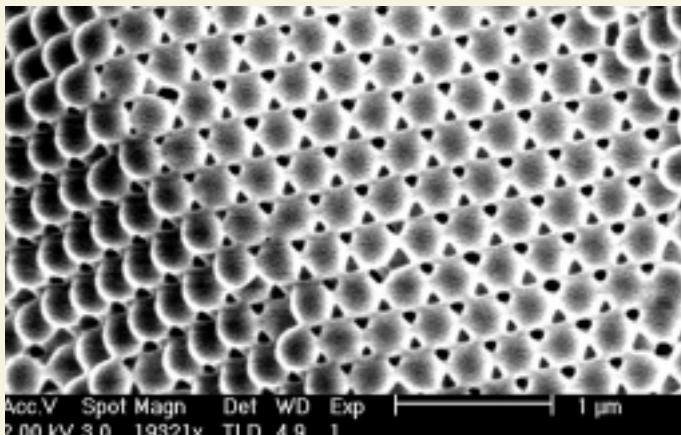
- F.U.R. (Forschungsgemeinschaft Urlaub und Reisen): Die Reiseanalyse RA 2003, Hamburg/Kiel 2003.
- ISENBERG, W./MÜLLENMEISTER, H. M./STEINECKE, A.: Tourismus im Wandel, Paderborn 2003.
- KULINAT, K.: Tourismusnachfrage: Motive und Theorien. – IN: BECKER, CHR./HOPFINGER H./STEINECKE, A. (Hrsg.): Geographie der Freizeit und des Tourismus - Bilanz und Ausblick, München/Wien 2003, S. 97-111.
- STEINECKE, A.: Tourismus und neue Konsumkultur: Kundenbedürfnisse – Schauplätze – Werthaltungen. – In: SCHNELL, P./POTTHOFF, K. E. (Hrsg.): Wirtschaftsfaktor Tourismus, Münster 2000, S. 81-91 (Münstersche Geographische Arbeiten, H. 42).
- WTO (World Tourism Organization): Tourism Highlights 2002, Madrid 2002.

Neue Gesetze des Lichts

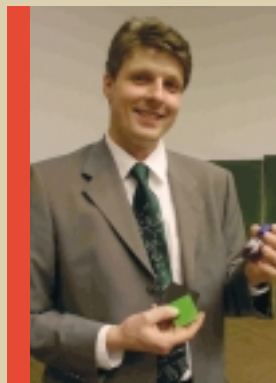
Die Idee der photonischen Kristalle

Die drei Aufgaben scheinen eines Harry Potters würdig: Bremse einen Lichtstrahl ab, bis seine Geschwindigkeit Null ist! Knicke, biege, teile ihn, manipulierte ihn nach allen Regeln der Kunst! Und dann – als Höhepunkt der Zauberei – nimm eine hell strahlende Lichtquelle und baue einen durchsichtigen Käfig so um sie herum, dass sie sich nicht mehr traut, auch nur ein einziges Photon auszusenden!

Ob die Magier aus Hogwarts die Aufgaben lösen könnten, ist schwer zu sagen. Physiker aus aller Welt haben sich jedoch genau dies vorgenommen: Die Gesetze des Lichts, wie sie bisher in der Schule gelehrt wurden, umzuschreiben oder zumindest Räume zu schaffen, wo sie neu interpretiert werden müssen. „Wir wollen Licht in Zukunft genauso beeinflussen können, wie es die Halbleitertechniker in den letzten Jahrzehnten mit den Elektronen geschafft haben“, sagt Ralf Wehrspohn.



Stapel Orangen – Kolloidaler Kristall.



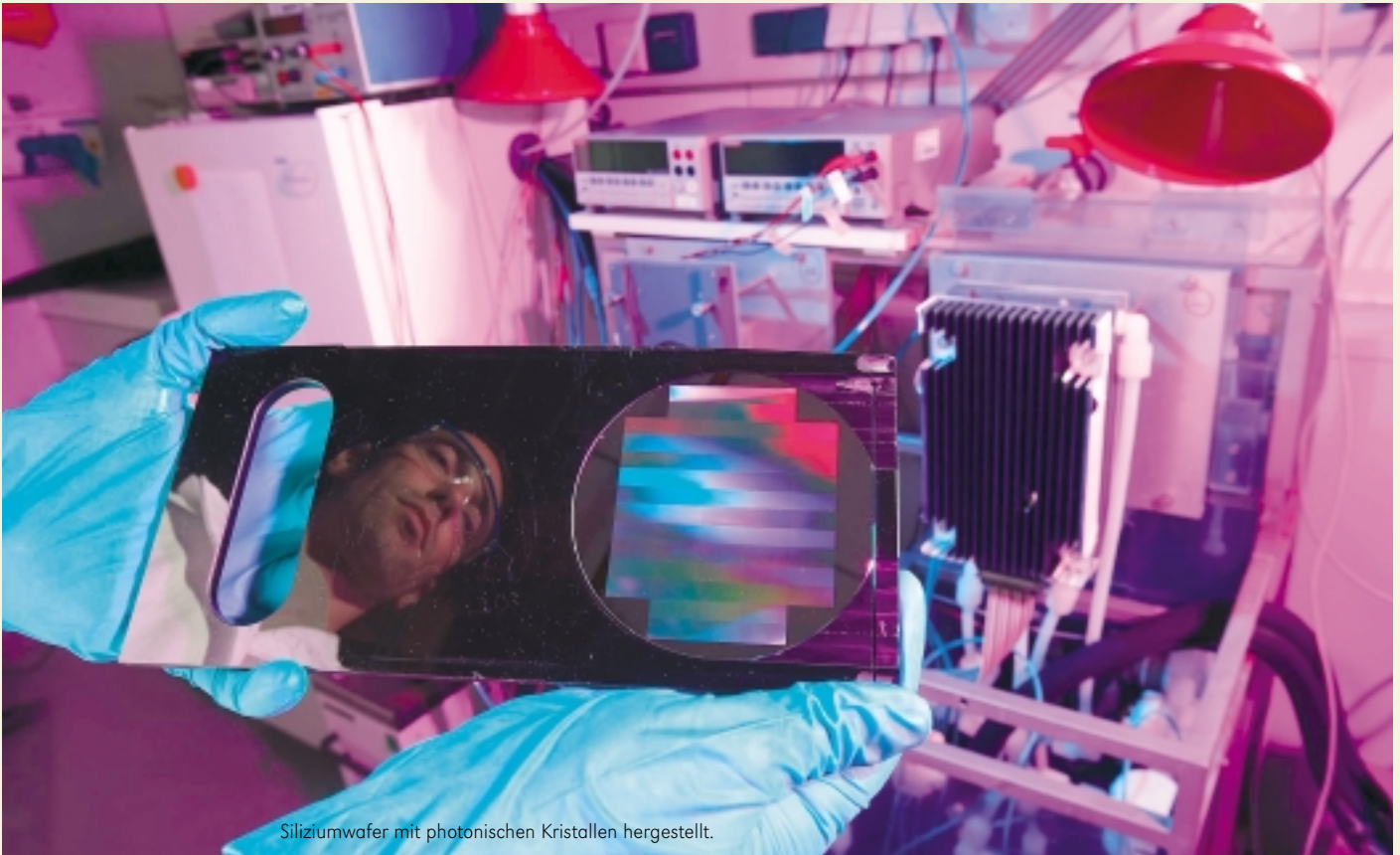
Prof. Dr. rer. nat. Dipl.-Phys.

Ralf B. Wehrspohn ist seit 2003 Professor für Experimentalphysik in der Fakultät für Naturwissenschaften an der Universität Paderborn. Seine Interessen liegen in der Simulation und Herstellung von Materialien für Nanophotonik wie z. B. photonische Kristalle oder Metamaterialien.

Die Idee der photonischen Kristalle, so heißt die Forschungsrichtung, ist noch sehr jung. Erst seit einigen Jahren gibt es erste kommerzielle Produkte im Bereich von Photonischen-Kristall-Glasfasern. Doch die Zukunftsaussichten sind grandios und reichen laut Wehrspohn „von der Sensorik über die Computer- und Telekommunikationsindustrie bis zur Kosmetikbranche“. Im Kern geht es bei den photonischen Kristallen um die Beeinflussung von Lichteigenschaften durch Interferenz – also um die Überlagerung von Lichtwellen. Diese Effekte kennt jeder, der schon einmal beobachtet hat, wie Wasserwellen ineinander fließen: Wo sich Wellenberge und -täler treffen, verstärken sie sich, wo hingegen ein Wellenberg auf ein Wellental trifft, löschen sich die beiden gegenseitig aus. Dasselbe gilt für Lichtwellen. Auch dies ist aus der Natur bekannt: So schimmert ein Ölfilm auf Wasser in unterschiedlichen Farben, je nachdem, unter welchem Winkel man ihn betrachtet. Dabei verstärken oder schwächen sich die Farben – die Wellen unterschiedlicher Wellenlänge –, die an der Vorder- und Rückseite des Films reflektiert werden. Diese Reflexionen wiederum gibt es, weil hier Materialien mit verschiedenen Brechungsindizes, nämlich Luft-Öl und Öl-Wasser, aufeinander stoßen. In der Technik nutzt man diesen Effekt ebenfalls schon seit längerem, beispielsweise für Antireflexbeschichtungen bei Brillen.

Halbleiter für Licht

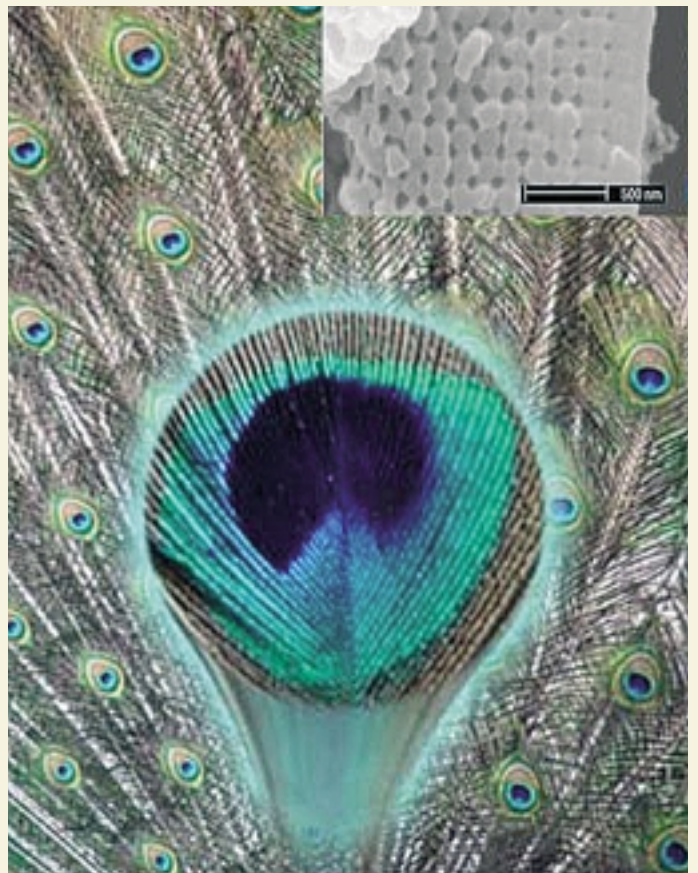
Ein dreidimensionaler photonischer Kristall ist nun einfach die logische Fortschreibung dieses physikalischen Prinzips in eine räumliche Form: Eine „periodische Änderung von Brechungsindizes in allen drei Raumrichtungen“, wie Wehrspohn sagt. Eli Yablonovitch, der als Forscher bei den Bell Labs in Holmdel, New Jersey, 1987 erstmals das Konzept der photonischen Kristalle



Siliziumwafer mit photonischen Kristallen hergestellt.

le beschrieb, stellte sich zur Illustration eine Kiste Orangen vor: Sitzt der Beobachter auf einer Orange in der Mitte, sieht er in allen Raumrichtungen dasselbe: Eine Abfolge von Luft – Orange – Luft – Orange und so weiter, eben eine periodische Änderung mit dem Orangendurchmesser als Periode. Wenn es nun in der Mikrowelt einen ähnlichen Aufbau gäbe, mit einer Periode in der Größenordnung von Lichtwellenlängen – also sozusagen Nano-Orangen mit Abmessungen von einigen hundert Nanometern –, dann, so überlegte Yablonovitch, könnten sich bestimmte Farben des Lichts in einem solchen Kristall überhaupt nicht ausbreiten. Der Grund: Die Reflexionen aus allen Raumrichtungen würden sich so überlagern, dass es zu einer vollständigen Auslöschung käme. In der Sprache der Physiker besitzt ein solcher Kristall für diese Lichtwellenlängen eine „vollständige Bandlücke“. Dies ist derselbe Effekt, den Halbleiter wie Silizium und Galliumarsenid für Elektronen aufweisen. Hier sind den Elektronen bestimmte Energieniveaus verboten – wegen dieser Analogie nennen die Physiker die photonischen Kristalle auch „Halbleiter für Licht“. Die Parallelen zur Halbleitertechnik sind so groß, dass 1987 ein weiterer Forscher, Sajeev John von der Universität Toronto, völlig unabhängig von Yablonovitch auf dieselbe Idee kam, indem er sich überlegte, ob sich bestimmte Konzepte aus der Elektronik auf Licht übertragen lassen. „Heute weiß man zwar durch zahlreiche theoretische Simulationen und Experimente, dass die Verhältnisse in den photonischen Kristallen komplizierter sind, als es das einfache Orangenmodell beschreibt, aber das Grundkonzept von Yablonovitch und John war völlig richtig“, sagt Wehrspohn.

andere, dass es sie in der Natur bereits gibt: Opale, Pfauenfedern, einige Fische und Insekten und bestimmte Schmetterlingsflügel enthalten Strukturen mit sehr regelmäßigen Löchern, die bewirken, dass diese Materialien, obwohl sie eigentlich transparent wären, in Rot, Grün oder Blau schimmern. Eben genau in



Pfauenfeder

Künstliche Opale

Während die einen Forscher Ende der 80er Jahre noch überlegten, wie sie photonische Kristalle herstellen könnten, entdeckten



3D-photonischer Kristall aus Silizium.

den Farben, die sich im Inneren nicht ausbreiten können und daher reflektiert werden. Mit dem Opal als Vorbild lassen sich inzwischen photonische Kristalle nach einer relativ simplen Methode herstellen: Winzige Quarz- oder Kunststoffkügelchen mit Durchmessern von einigen hundert Nanometern werden in einer Flüssigkeit gelöst, die dann so langsam verdunstet, dass sich die Kügelchen von selbst zu den „Nano-Orangen-Stapeln“ anordnen – und schon hat man einen photonischen Kristall, der in verschiedenen Farben schillert. „Derartige künstliche Opale werden beispielsweise demnächst als Ersatz für Lebensmittel- oder Kosmetikfarbstoffe auf den Markt gebracht“, weiß Siegmund Greulich-Weber. „Der Vorteil gegenüber herkömmlichen Pigmenten ist, dass die neuen Farbstoffe nur aus Quarzteilchen, also Sand, bestehen und damit völlig unbedenklich sind.“ Einen noch besseren photonischen Kristall, einen so genannten invertierten Opal, erhalten die Forscher, wenn sie ein Material wie Zinnsulfid oder Titanoxid in die Hohlräume zwischen den Kugeln füllen und dann die Kugeln durch Hitze oder Ätzverfahren entfernen. Übrig bleibt eine Wabenstruktur, eben der invertierte Opal, mit einer vollständigen Bandlücke. Neben der Arbeitsgruppe von Heinz Kitzerow, in der abstimmbare photonische Kristalle auf der Basis invertierter Opale hergestellt werden, arbeiten auch Siegmund Greulich-Weber und Mitarbeiter daran für z. B. Anwendungen im Bereich der Photovoltaik. So können künstliche Opale großflächig auf Solarzellen aus spektral-selektive und räumlich-beugende Filter aufgebracht werden, um die Effizienz von Stapelsolarzellen zu erhöhen. Diese Arbeiten erfolgen in sehr enger Zusammenarbeit der Kollegen aus der Chemie (Prof. Heinrich Marsmann, Prof. Claudia Schmidt, und Prof. Klaus Huber).

Da viele nützliche Anwendungen aber bereits mit zweidimensionalen photonischen Kristallen funktionieren – die also nur in zwei Richtungen eine periodische Variation des Brechungsindex aufweisen –, verwendet die Mehrzahl der heutigen Forscher die Methoden, die am Max-Planck-Institut für Mikrostrukturphysik in Halle aus der ehemaligen Arbeitsgruppe von Ralf Wehrspohn perfektioniert worden sind: Sie ätzen mit ausgeklügelten Verfahren eine Vielzahl winziger Kanäle in Materialien wie Silizium. Diese Poren sind extrem regelmäßig angeordnet, haben Durchmesser und Abstände von einigen hundert Nanometern und können hundertmal tiefer sein, als sie breit sind. Derartige Bauteile sind ideale photonische Kristalle für Infrarotlicht, wie man es beispielsweise zur Datenübertragung in der Telekommunikation braucht. Silizium ist im Infraroten normalerweise transparent, aber nicht wenn dasselbe Licht senkrecht zu den Poren durch den Kristall geleitet wird. Dann wirkt der Kristall für bestimmte Wellenlängen, die etwa der doppelten Periodenlänge entsprechen, wie ein perfekter Spiegel. Sie können sich in ihm nicht ausbreiten.

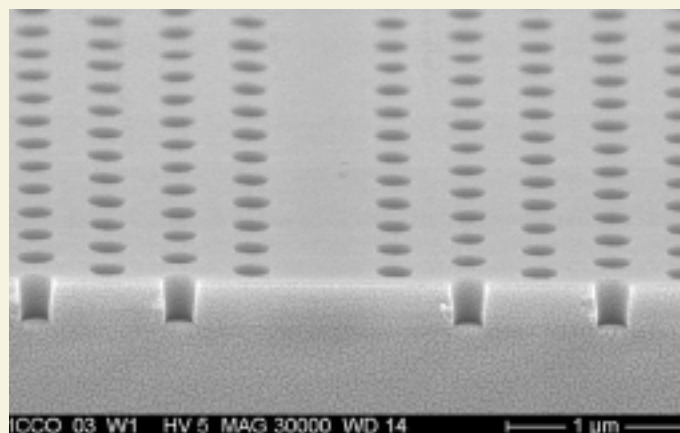
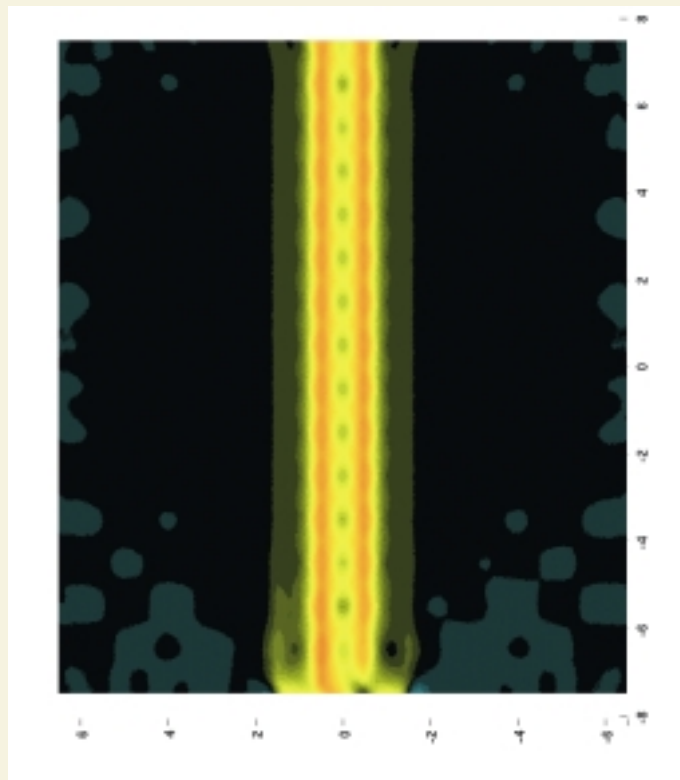
Photonische Kristalle

„In der heutigen Optik passt etwa ein Bauteil auf einen Kubikmillimeter“, erklärt Wehrspohn, „mit photonischen Kristallen werden es künftig tausendmal so viele sein“. Das erklärt auch das hohe Interesse, das Firmen an dieser Technik haben: So koordiniert Infineon derzeit ein vom Bundesforschungsministerium gefördertes Projekt, das sich zum Ziel gesetzt hat, die nötigen Komponenten für eine künftige optische Datenverarbeitung zu entwickeln.

„Die Idee dabei ist, künftig Daten per Glasfaser bis in die Haushalte zu transportieren“, sagt Ralf Wehrspohn. Datenraten von 1 000 Gigabit pro Sekunde bis ins heimische Büro – zehntausendmal mehr als mit ISDN – wären mit Glasfasern machbar, „aber nur, wenn die dazu nötige optische Datenverarbeitung kompakt und kostengünstig auf einen Chip passt, an den dann einige hundert Haushalte angeschlossen sind. Heute sind derartige Geräte extrem teuer und füllen ganze Räume“. Die photonischen Kristalle sind dafür wie geschaffen: Sie können helfen, die Lichtstrahlen in die unterschiedlichen Frequenzen aufzuteilen, die informationstragenden Impulse zu extrahieren oder zu addieren, zu verstärken und weiterzuleiten und das alles extrem kompakt und auf demselben Chipmaterial.

Denn erst die photonischen Kristalle bieten den Forschern eine Vielzahl von Möglichkeiten, Licht nach Belieben zu manipulieren:

- Das Grundmaterial sowie die Durchmesser, Abstände und Anordnungen der Löcher legen für jede Wellenlänge des Lichts den im Kristall gültigen Brechungsindex fest. „Da die Lichtgeschwindigkeit im Kristall gleich der Vakuumlichtgeschwindigkeit geteilt durch den Brechungsindex ist, können wir damit tatsächlich Licht bestimmter Wellenlängen beliebig verlangsamen oder zum Stillstand bringen“, sagt Prof. Wehrspohn.
- Darüber hinaus können die Physiker ebenso wie in den Halbleitern der Elektronik gezielt „Defekte“ in den Kristall einbringen, etwa Löcher weglassen oder Straßen und Weggabelungen – Wellenleiter – konstruieren. An diesen Orten kann dann Licht, das sich normalerweise in diesem Kristall nicht ausbreiten darf, plötzlich doch existieren und das in extrem hoher Intensität. Wenn – was allerdings rechenintensive Simulationen erfordert – die Anordnung der Defekte richtig gewählt wird, kann Licht zum Beispiel wesentlich schärfer um Kurven geführt werden, als es mit anderen Materialien jemals möglich wäre.
- Setzt man eine Lichtquelle, die Licht einer im Kristall verbotenen Wellenlänge emittiert, in einen photonischen Kristall, dann wird kein einziges derartiges Photon emittiert. So lässt sich die spontane Emission von Lasern unterdrücken, die deren Hauptverlustquelle ist. Auf diese Weise könnten wesentlich effizientere Laser gebaut werden.
- Die Löcher der photonischen Kristalle müssen nicht unbedingt mit Luft gefüllt sein. „Befänden sich etwa Flüssigkristalle darin, wie Sie in Flachbildschirmen verwendet werden, könnten wir während des Betriebs den Brechungsindex und damit die Eigenschaften des Materials ändern, indem wir eine elektrische Spannung oder ein Magnetfeld anlegen“, sagt Heinz Kitzerow aus der Physikalischen Chemie.
- Doch auch für die Sensorik eignen sich die Kristalle gut: Wehrspohns Team arbeitet beispielsweise zur Zeit mit den Firmen Dräger, Infineon und dem Fraunhofer-Institut für Physikalische Messtechnik in Freiburg an einem Sensor, der Alkohol in der Atemluft messen kann: Der Alkohol wird dabei durch die Löcher des photonischen Kristalls geblasen, wobei die winzige Absorptionsänderung von 1 zu 200 Milliarden ausreicht, um aussagekräftige Resultate zu bekommen. Derartige Sensoren werden in den USA als Drive-Interlocks für Fahrzeuge gefordert: Sie sollen verhindern, dass das Auto



Wellenleiter in photonischen Kristallen, Simulation dazu.

angelassen werden kann, wenn der Fahrer eine „Alkoholfahne“ hat. Heutige Systeme sind mindestens handtellergröße Kästen, mit photonischen Kristallen könnten die Geräte kleiner als eine Fingerkuppe sein.

Anders als bei vielen neuen Entdeckungen, die oft jahrelang nur Grundlagenforscher interessieren, haben bei den photonischen Kristallen Firmen schnell die Potenziale erkannt und versuchen, sie zu nutzen. Einige Produkte – etwa Fasern mit photonischen Kristallen, die weißes Licht erzeugen, welches zehntausendmal heller ist als das Sonnenlicht – gibt es bereits zu kaufen und werden in Paderborn in der Arbeitsgruppe von Artur Zrenner zur Spektroskopie eingesetzt.

Optische Interconnects

Ein weiteres Entwicklungsgebiet sind optische Interconnects. Das sind optische Datenverbindungen zwischen den Prozessoren und Speicherchips eines Computers, die einen der entscheidenden Engpässe der Computertechnik beheben könnten: Sie könnten die elektrischen Leitungen ersetzen, über die die Daten zehn-



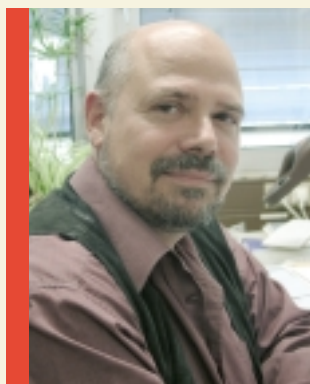
Zum CeOPP (Center for Optoelectronics and Photonics) gehören derzeit zehn Arbeitsgruppen aus den Departments Chemie, Elektrotechnik und Physik.

mal langsamer fließen als innerhalb der Mikrochips. Diese Busleitungen begrenzen die Leistungsfähigkeit heutiger Computer – weitere Verbesserungen der Prozessoren bringen nichts, wenn es nicht gelingt, den Datenaustausch zwischen den Chips zu beschleunigen und besser zu isolieren. Bestünden die Platinen aus photonischen Kristallen, könnte der Datenaustausch optisch erfolgen, was zudem den Vorteil hätte, dass die Mikrochips näher aneinander rücken könnten, weil optische Leitungen, anders als elektrische, sich nicht gegenseitig stören. In diesen Geräten würden dann immer noch die Elektronen dazu dienen, die Rechenoperationen durchzuführen, aber mehr und mehr würde die Datenübertragung mit Licht erfolgen. Sehr wichtige theoretische Entwicklungen zu optischen Interconnects sind in Paderborn in der Arbeitsgruppe von Prof. Gerd Mroczynski aus der Elektrotechnik gelaufen.

Der

Optische Transistor

„In fernerer Zukunft“, sagt Wehrspohn, „liegt das, was für viele Physiker der Heilige Gral der Photonik ist: eine rein optische Datenverarbeitung, in der die Photonen auch die Rechnungen durchführen – was die Computer hundert- bis tausendfach schneller machen würde. Eine wichtige Komponente dafür könnte vielleicht in zehn bis fünfzehn Jahren funktionieren: Der optische Transistor. Dafür braucht man dann einen perfekten dreidimensionalen photonischen Kristall mit vollständiger Bandlücke. In einem solchen Material ist es möglich, mit einem Lichtimpuls die Durchlässigkeit des Stoffes für einen anderen Lichtstrahl zu verändern – dann schalten wir tatsächlich Licht mit Licht“, sagt Ralf Wehrspohn. Das wäre dann endgültig der Beginn des photonischen Zeitalters. Das dieses in Paderborn schon begonnen hat, spiegelt sich an dem interdisziplinären „Center für Optoelectronics und Photonics“ (CeOPP) wieder, an dem Physiker, Chemiker und Elektrotechniker seit 2004 gemeinsam forschen und lehren.



apl. Prof. Dr. rer. nat Dipl.-Phys.

Siegmund Greulich Weber

ist seit 2004 außerplanmäßiger Professor in der Experimentalphysik. Seine Arbeitsgebiete sind Festkörperphysik, Festkörperspektroskopie, Magnetische Resonanzspektroskopie mit Schwerpunkten in Siliziumkarbid und Photonische Kolloidkristalle.



Paderborns gute Adresse

Marienplatz 2a • 33098 Paderborn • Tel.: 0 52 51/ 88-29 80
Fax: 88-29 90 • tourist-info@paderborn.de • www.paderborn.de

Öffnungszeiten:

(April-Oktober)

Mo.-Fr. 10-18 u. Sa. 10-14 Uhr

(November-März)

Mo.-Fr. 10-17 u. 10-14 Uhr

Tourist

Information Paderborn



Grid Computing

Kooperation und gemeinsame Nutzung vernetzter Ressourcen

Das Internet ist ein weltumspannendes Netzwerk, an das eine immense Zahl von Rechnersystemen angeschlossen ist. Grid-Computing ermöglicht die Zusammenfassung all dieser Systeme zu einem virtuellen Supercomputer, der jedem Teilnehmer für seine eigenen Anwendungen zur Verfügung steht. Ähnlich der Stromversorgung in einem Haushalt kann ein Benutzer bequem die Rechenleistung dieses virtuellen Supercomputers abrufen und seinen Verbrauch bezahlen. Ein Softwaresystem, genannt Grid-Middleware, sorgt für Benutzerfreundlichkeit, indem es dem Benutzer einen bequemen, intuitiven Grid-Zugang bietet und die Komplexität des Gesamtsystems verbirgt. Ebenso wie das World Wide Web grundlegend die Art und Weise verändert hat, wie wir uns informieren, wird das Grid-Computing die Nutzung von Computer-Ressourcen revolutionieren. Dieser Artikel beschreibt die Grid-basierte Forschung, die im Paderborn Center for Parallel Computing (PC²) und am Institut für Informatik der Universität Paderborn durchgeführt wird.

Für Außenstehende mag es verwunderlich klingen, wenn wissenschaftliche und kommerzielle Einrichtungen stets aufs Neue in teure Hochleistungsressourcen investieren, etwa in schnelle Supercomputer mit Tausenden von Prozessoren, für deren Unterbringung zum Teil Gebäude in der Größe einer Mehrzweckhalle errichtet werden. Führt man sich jedoch vor Augen, dass ein aktueller PC mehr als zwei Wochen an einer einzigen Simulation eines Crashtests rechnen würde, so wird die Notwendigkeit von Hochleistungsressourcen zur Verkürzung der Entwicklungsprozesse oder zum Erlangen von präzisen und umfassenden Erkenntnissen deutlich.

Die benötigten Hochleistungsressourcen werden in der Regel als Parallelrechner realisiert. Anfangs waren dies speziell angefertigte Maschinen mit besonderen Hardware- und Softwarekomponenten, deren Anschaffung sich lediglich große Rechenzentren und Konzerne leisten konnten. Danach wurden immer mehr Elemente dieser Architekturen durch Standardkomponenten ersetzt. Seit Mitte der 90er Jahre dominiert die Clustertechnologie, die durch Vernetzung und koordinierte Verwendung vieler Standard-PCs die Leistungsfähigkeit großer Parallelrechner annähernd erreicht. Damit wurde Hochleistungsrechnen für eine breite Schicht von Anwendern erschwinglich, die sich durch die Installation eines Clusterrechners eine vertrauenswürdige Rechenumgebung zur eigenen exklusiven Nutzung schufen.

Diese Strategie der Ressourcennutzung ist allerdings mit einigen Nachteilen behaftet. Zu den Anschaffungskosten addieren sich noch erhebliche Kosten für die Administration, Wartung und

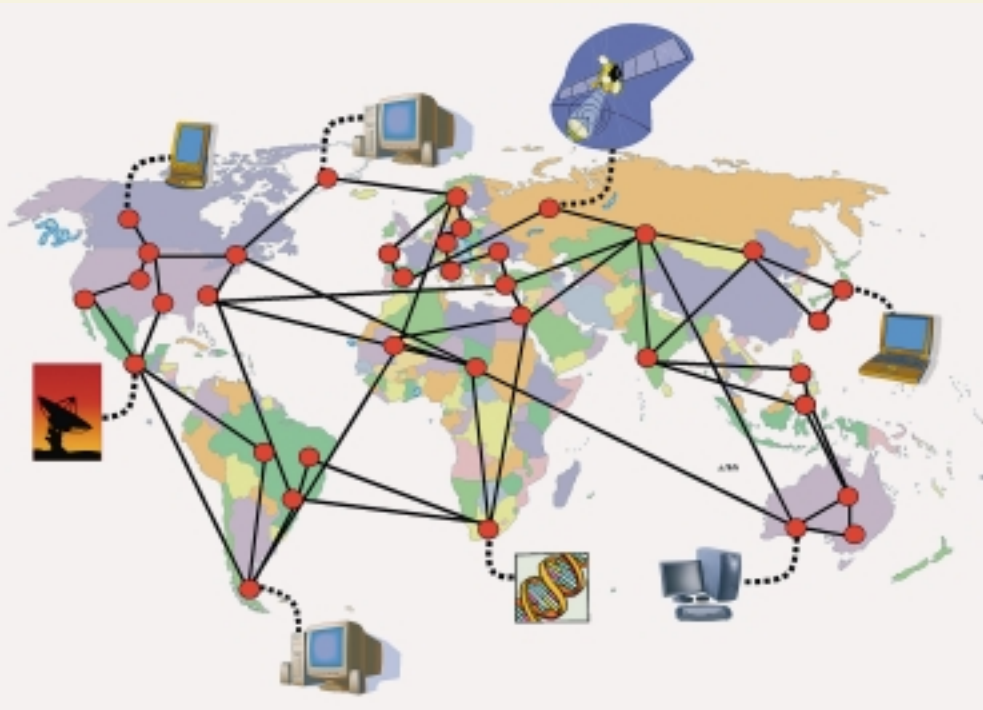


Prof. Dr. rer. nat. Odej Kao leitet die Arbeitsgruppe „Betriebssysteme und Verteilte Systeme“ am Institut für Informatik und ist geschäftsführender Leiter des Paderborn Center for Parallel Computing (PC²). Seine Forschungsinteressen umfassen Ressourcenmanagement für Grid- und Cluster-Computing, parallele und verteilte Multimediainformationssysteme sowie Kooperation in Wireless LAN-Umgebungen.

spätere Erweiterung der Hard- und Software. Ferner zeigen Erfahrungswerte, dass über einen längeren Zeitraum betrachtet die Auslastung der Ressourcen niedriger ist als erwartet. Nichtsdestotrotz sind zu Stosszeiten so viele Aufträge zu bearbeiten, dass die Ressourcen überlastet sind und Wartezeiten entstehen. Hierdurch wird der gesamte Entwicklungs- und Produktionsprozess verzögert. Aus diesen Gründen werden Ressourcen von vielen Anwendern nicht mehr gekauft, sondern bei Rechenzentren oder anderen professionellen Rechnerbetreibern angemietet. Der Wegfall von Administrations- und Wartungskosten geht jedoch mit dem Verlust des exklusiven Nutzungsrechts einher, so dass Wartezeiten bei der Auftragsannahme und -bearbeitung auftreten können.

Grid-Computing

Neben Kauf und Miete bietet das Grid-Computing eine dritte Möglichkeit, den Bedarf an Rechen-, Speicher- und anderen Ressourcen zu decken. Die Grundlage bilden die bereits weltweit existierenden Rechen- und Speicherkapazitäten, die mittels Internet gekoppelt sind und daher von jedem Netzzugang aus erreicht werden können. Ein Anwender, dessen Ressourcen aktuell ungenutzt sind, kann diese anderen Nutzern zur Miete anbieten. Auf der anderen Seite können Nutzer, die aktuell mehr Rechenbedarf haben als die eigenen Ressourcen zu leisten vermögen, die angebotenen Ressourcen in Anspruch nehmen und die verbrauchte Rechenzeit bezahlen. Damit diese Interaktion erfolgreich und weltweit umgesetzt werden kann, wird eine spezielle Software benötigt, die so genannte *Grid-Middleware*, die alle Angebote und Anfragen registriert, die technischen Eigenschaften der Ressourcen verwaltet, die Zuordnung von Aufträgen zu Ressourcen übernimmt und schließlich die Komplexität dieser Vorgehensweise vollständig hinter einer



komfortablen Schnittstelle vor dem Benutzer versteckt. Dieses Prinzip ist uns aus dem Alltagsleben bestens bekannt. Ein Haushalt würde kaum seinen Strombedarf vollständig durch ein eigenes Kraftwerk decken, sondern bezieht die aktuell benötigte Menge Energie über Leitungen von einem Stromversorger. Gleiches gilt für die Versorgung mit Wasser oder für den Zugang zu Informationen unter Verwendung von Telefonleitungen und Internet Providern. Für die Waren *Strom, Wasser und Information* existieren somit Infrastrukturen, die einen komfortablen und bedarfsorientierten Zugang zur jeweiligen Ware ermöglichen und ein mengenbezogenes Abrechnungssystem zugrunde legen. Die Übertragung dieser Vorgehensweise auf den Rechenzeitbedarf war und ist Gegenstand vieler Projekte in der Informatik. So wurde bereits in den 60er Jahren das Rechnersystem MULTICS (*Multiplexed Information and Computing Service*) entwickelt, das ausreichend Rechenkapazität für alle Bewohner des Ballungszentrums Boston zur Verfügung stellen sollte [1]. Ein prominentes Beispiel der neueren Zeit ist das *Application Service Providing (ASP)*, bei dem Benutzer bei Bedarf Zugang zu ausgewählten Standardanwendungen erhalten und pro Nutzungszeit bezahlen. Die Anwendbarkeit dieser Idee wird auch in vielen anderen Bereichen der Informationstechnologie unter Beweis gestellt:

- Vermietung von Speicherplatz (*Storage Service Providing*), beispielsweise für den Internetauftritt oder die redundante Datenhaltung
- Internet- bzw. Peer-to-Peer-Computing: Zusammenfassung vieler Rechner im Internet zur Bearbeitung von rechenintensiven, aber nicht zeitkritischen Problemen, wie zum Beispiel der Suche nach außerirdischer Intelligenz (*seti@home*) oder der Entschlüsselung des menschlichen Genoms (*Genome@home*).
- Webdienste wie etwa Suchmaschinen, die einen komfortablen Zugang zu spezialisierten Anwendungen und den zugeordneten Rechenressourcen bieten.

Wodurch unterscheidet sich Grid-Computing von diesen Ansätzen? Zunächst wird bei Grid-Computing berücksichtigt, dass Organisationen wie Universitäten, Forschungszentren, Behörden

und Betriebe bereits Hochleistungsressourcen betreiben und diese weiter nutzen wollen. Daher werden keine signifikanten Änderungen oder Erweiterungen der bestehenden Infrastruktur vorausgesetzt, lediglich die Installation der Grid-Middleware ist notwendig. Ferner erfolgt die Interaktion ohne eine zentrale Anmelde- und Verwaltungsstelle. Die Teilnahme an der Grid-Infrastruktur wird dynamisch und bei Bedarf an- und abgemeldet. Schließlich bietet das Grid-Computing unmittelbare Vorteile für Betreiber und Anwender von Ressourcen. Ein Betreiber kann wertvolle Rechenressourcen verwenden, die insbesondere zu ungünstigen Zeiten (z. B. nachts, an Wochenenden oder zu Ferienzeiten) nicht durch Aufträge eigener Benut-

zer ausgelastet sind und somit einen Teil der Investitionskosten zurückholen. Bei der Neuanschaffung können die Systeme gemäß der gewünschten Antwortzeiten im täglichen Betrieb dimensioniert werden. Eine Überlastung zu Stosszeiten wird durch Verlagerung von Aufträgen in die Grid-Umgebung abgefangen, so dass die Ergebnisse trotz einer kleineren Maschine in akzeptabler Zeit vorliegen. Gerade für Unternehmen bedeutet dies einen unschätzbaren Vorteil hinsichtlich ihrer Flexibilität: Aufträge können jederzeit angenommen und kurzfristig erledigt werden, da nicht erst eigene Rechenressourcen aufgebaut werden müssen. Außerdem ermöglicht Grid-Computing die Erprobung verschiedener Rechnerplattformen bis hin zu seltener und teurer Hardware, bevor langfristige strategische Entscheidungen getroffen werden.

Grundkomponenten des Grid-Computing

Häufig verwendete Definitionen charakterisieren das Grid-Computing als eine flexible, sichere und koordinierte Ressourcennutzung in einer dynamischen Gruppe von Individuen, Institutionen und Ressourcen zum Erreichen eines gemeinsamen Ziels [2].

Das erste wesentliche Element des Grid-Computing sind demnach Ressourcen. Darunter werden alle Prozessoren, Rechner, Speichermedien, Netzwerke und alle anderen Bestandteile einer Rechnerarchitektur verstanden. Aber auch spezielle Ein- und Ausgabegeräte wie Teleskope in der Astronomie, Detektoren in der Atomphysik oder hochauflösende 3D-Displays in der Visualisierung gehören dazu und können kooperativ genutzt werden. Neben diesen Hardware-Ressourcen müssen für die Auftragsausführung im Grid auch passende Software-Ressourcen bereitgestellt werden. Hierzu zählt das Betriebssystem ebenso wie die verfügbaren Anwendungen samt benötigter Lizenzen und Laufzeitbibliotheken. Der Begriff Ressource umfasst im Grid-Kontext auch Informationen, Datenbestände und Wissen, beispielsweise Archive mit astronomischen Daten oder Informationen über das menschliche Genom. Und schließlich zählen sogar Menschen zu den Ressourcen. Beispielsweise kann ein

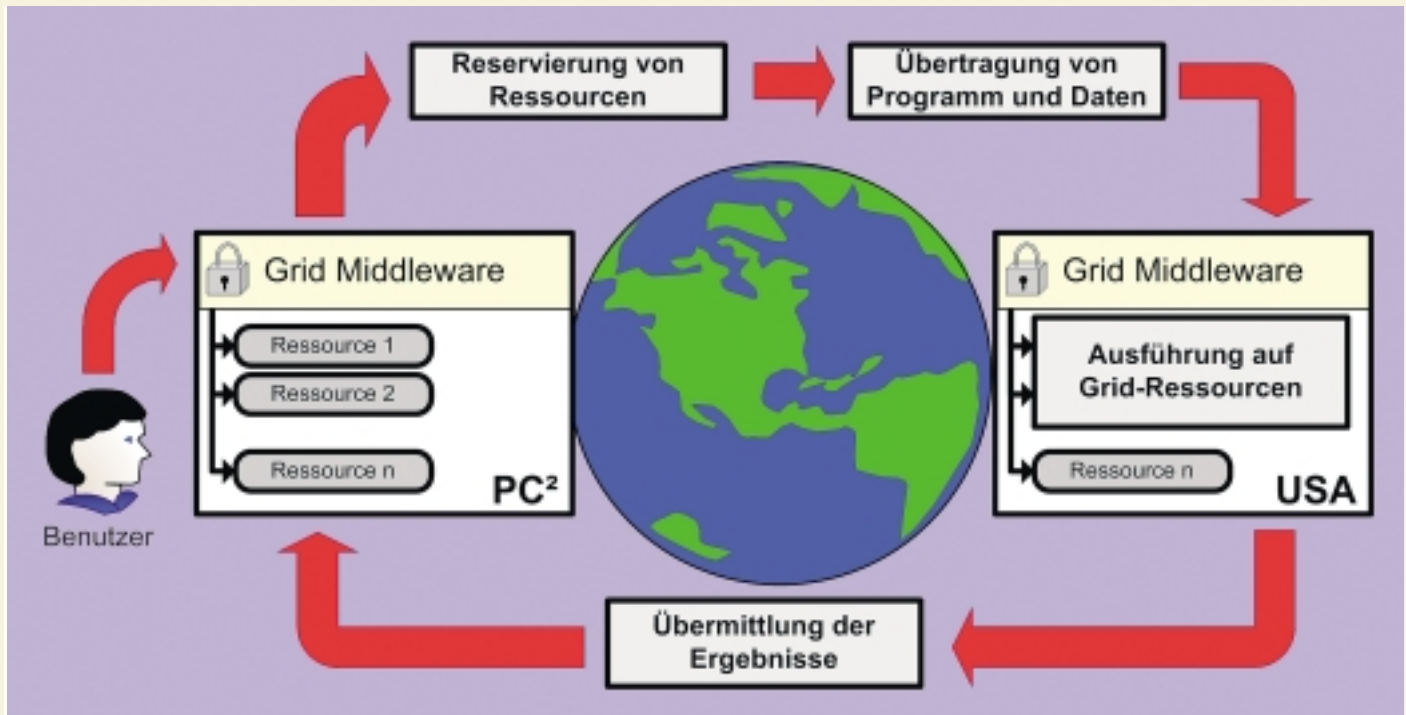


Abb. 1: Grundsätzlicher Ablauf bei Ausführung von Aufträgen in einer Grid-Umgebung.

Anwender verlangen, dass sein Auftrag nur an solche Grid-Betreiber vergeben werden darf, bei denen ein Experte für eine bestimmte Anwendung arbeitet und der notfalls über die Support-Hotline erreichbar ist.

Die Komplexität und die Vielfalt der eingebundenen Ressourcen wird deutlich, wenn man sich die Größe des heutigen Internets und die Vielzahl unterschiedlicher Betriebssysteme, Anwendungen oder Compiler vor Augen führt. Die bequeme, zuverlässige, sichere und für den Anwender transparente Nutzung der Grid-Infrastruktur ist daher die zentrale Aufgabe der Grid-Middleware. Die Software verwaltet den Grid-Zugriff auf die Ressourcen einer administrativen Einheit wie etwa eines universitären Rechenzentrums und kooperiert mit der Grid-Middleware der anderen Grid-Teilnehmer. Dazu werden zunächst statische Informationen (Hardwaretyp, Betriebssystemversion, eingesetzte Netzwerke, ...) und dynamische Informationen (aktuelle Belegung, wartende Aufträge, Ausfälle, ...) über die Ressourcen gesammelt und verwaltet. Diese Informationen werden an diejenigen Grid-Teilnehmer übermittelt, welche die lokalen Ressourcen für ihre Berechnungen nutzen wollen. Auf der anderen Seite kann die Middleware Informationen über alle anderen Ressourcen im Grid zusammentragen, die von den eigenen Benutzern benötigt werden. Diese Informationsbeschaffung und Präsentation erfolgt über eine einheitliche und bequeme Schnittstelle, so dass die Komplexität des Grids verborgen bleibt: Der Anwender muss nicht wissen, wie er Zugang zu den Grid-Ressourcen erhält oder wie Programme dort zu starten sind. Ihm muss nicht einmal bewusst sein, ob seine Aufträge auf den lokalen oder auf entfernten Ressourcen im Grid ausgeführt werden. Die Grid-Middleware lässt also die administrativen Grenzen zwischen den Ressourcen verschwinden, wobei man von der Virtualisierung aller Ressourcen spricht.

Es stellt sich hierbei die Frage, woher die Grid-Middleware weiß, welche Ressourcen aktuell benötigt werden. Dies wird vom Benutzer bei der Übergabe des Auftrags mitgeliefert, zum Beispiel dass ein Cluster mit Infiniband-Netzwerk benötigt wird,

dessen Knoten über Pentium 4 Prozessoren und mindestens 4 GByte Hauptspeicher verfügen, und auf denen das Betriebssystem Linux mit Kernel 2.6 installiert ist. Grundsätzlich kann der Benutzer jede Ressource anfordern, die im Grid-System verfügbar ist. Eine analoge Vorgehensweise wird beispielsweise auch bei einer Flugbuchung angewendet (Start/Ziel, Zeit, Klasse, Personen, Fluggesellschaft, ...).

Basierend auf diesen Angaben sucht die Middleware passende und verfügbare Grid-Ressourcen. Der Anwender bekommt eine Liste zurück, in der Angaben wie Startzeit und Preis der Nutzung vermerkt sind. Aus dieser Liste wählt der Benutzer ein Angebot und initiiert die Verarbeitung seines Auftrags. Die folgenden Schritte werden dabei automatisch von der Middleware ausgeführt:

1. Reservierung gefundener Ressourcen für den erteilten Auftrag (Allokation und Scheduling)
2. Übertragung des Programmcodes und der benötigten Daten zum entfernten Grid-System. Dabei werden die Daten verschlüsselt und der Absender auf Wunsch anonymisiert, so dass von Dritten keine Rückschlüsse auf die Aktivität gezogen werden können.
3. Der Auftrag wird auf den entfernten Grid-Ressourcen zur Ausführung gebracht. Die Middleware wartet derweil auf das Ende der Verarbeitung.
4. Sobald die Ergebnisse vorliegen, werden diese an den Anwender zurückgesendet und auf seinem eigenen Rechner gespeichert. Der Anwender kann diese Ergebnisse nun mit einer lokalen Anwendung oder einem Browser auswerten.

Die einzelnen Schritte dieses Ablaufs finden sich in Abbildung

1. Die bekanntesten Beispiele für existierende Middleware, die diesen Ablauf realisieren, sind Globus [3] und Unicore [4].

PC² Forschungsschwerpunkt:

Zuverlässigkeit der Auftragsausführung im Grid

Es ist deutlich geworden, dass sich in einem Grid-System zwei Gruppen gegenüber stehen: Auf der einen Seite die Anwender,

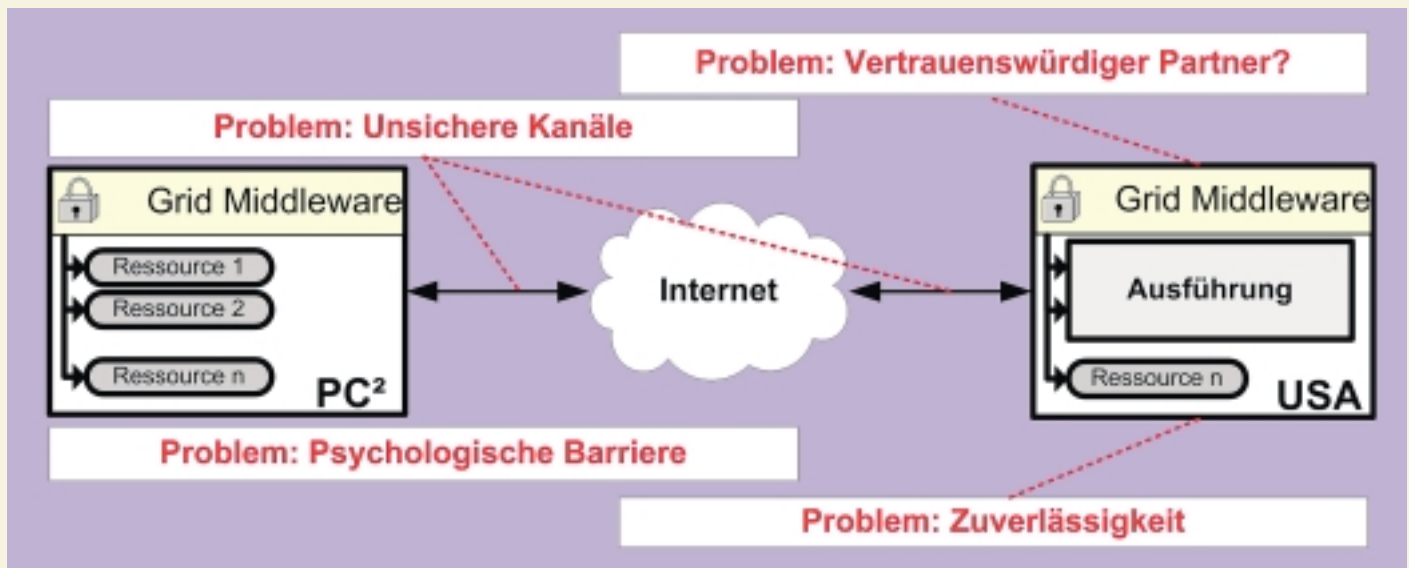


Abb. 2: Herausforderungen im Grid-Computing.

die das Grid für ihre Zwecke nutzen wollen und hierfür bestimmte Ressourcen benötigen; auf der anderen Seite die Provider, die ihre Grid-Ressourcen zu einem möglichst hohen Grad auslasten möchten. Um Akzeptanz zu finden, muss ein Grid-System diese beiden Gruppen derart zueinander führen, dass die Anforderungen jeder Seite erfüllt werden.

Um diese Anforderungen schriftlich zu fixieren, wird in einem ersten Schritt zunächst ein so genanntes „Service Level Agreement“ (SLA) zwischen dem Grid-Benutzer und den beteiligten Providern ausgehandelt. In diesen Verhandlungsprozess werden nicht alle Provider eingebunden, sondern nur diejenigen, welche die vom Benutzer geforderten Rahmendaten erfüllen können. Solch ein SLA beschreibt alle Erwartungen und Verpflichtungen in der Beziehung zwischen Grid-Nutzer und Provider und kann als vertragliche Grundlage der Zusammenarbeit angesehen werden.

Die auszuhandelnden Rahmendaten betreffen vor allem die folgenden Aspekte:

- **Ressourcen:** Beschreibung der verschiedenen Hardware-, Software-, Informations- und menschlichen Ressourcen, die für die Durchführung der Anwendung notwendig sind.
- **Dienstgüte und Fehlertoleranz:** Bei der Durchführung des Auftrags muss das System Prioritäten, Fristen oder Grenzwerte einhalten. Vor allem der kommerzielle Grid-Benutzer will sicher sein, dass er seine Berechnungsergebnisse zum vereinbarten Termin erhält. Einem Wetterdienst würde etwa die Berechnung des Wochenendwetters wenig nutzen, wenn sie wegen ausgefallener Grid-Ressourcen erst am Montag abgeschlossen werden würde.
- **Geheimhaltung:** Bei besonders sensiblen Berechnungen muss sicher gestellt sein, dass geeignete Schutzmaßnahmen getroffen werden. Diese reichen von der Geheimhaltung der Nutzung, über Firewalls zum Schutz vor Angriffen und unberechtigten Zugriffen, bis hin zu Sicherheitsanforderungen an das ausführende Rechenzentrum.
- **Kosten und Kontrolle:** Die Nutzungskosten für die Grid-Ressourcen hängen von den Vereinbarungen des SLAs ab. Je höher die Ansprüche an Geschwindigkeit, Dienstgüte oder Sicherheit, desto höher der Preis. Die Vereinbarung einer Vertragsstrafe regelt, wie viel der Provider dem Kunden zu

zahlen hat, falls die Leistungen des SLAs nicht erbracht werden.

Sind sich Grid-Benutzer und Provider über alle Konditionen einig, so ist der Verhandlungsprozess erfolgreich und zwischen ihnen wird ein SLA geschlossen. Der Kunde erhält somit die Garantie, dass seine Berechnung wie gewünscht durchgeführt wird. Der Provider wiederum kann die Nutzung durch den Grid-Benutzer in die Auslastung seiner Ressourcen einplanen [5].

Ein entscheidendes Kriterium für die Akzeptanz der SLAs ist die autonome Auswertung der eingegebenen Anforderungen und Mitteilung der Bearbeitungskonditionen seitens des Providers, so dass der Benutzer lediglich die Auswahl treffen muss. Die Bereitstellung einer entsprechenden Sprache für die Definition maschinenlesbarer SLAs für einzelne oder verkettete Aufträge (Workflows) ist aktuell ein wesentlicher Forschungsgegenstand.

Sobald ein SLA geschlossen wurde, muss der Provider geeignete Mechanismen zur Einhaltung der abgegebenen Garantie aktivieren. Dies bedeutet zunächst, dass der Provider die angeforderten Ressourcen reservieren und dem Kunden so zuordnen muss, dass der Auftrag vor der Sollzeit erledigt ist. Dies soll insbesondere dann der Fall sein, wenn einzelne Komponenten des Zielsystems ausfallen, d. h. das Resource-Management-System (RMS) muss auf Ausnahmesituationen reagieren. Vom Defekt eines Rechenknotens ist nicht nur die aktuell laufende Berechnung betroffen, sondern auch alle zukünftigen Berechnungen auf diesem Knoten. Das RMS muss folglich die Menge der noch zur Verfügung stehenden Ressourcen neu unter den Aufträgen aufteilen, so dass die Einhaltung abgegebener Garantien gesichert bleibt.

Dies vermögen jedoch existierende Resource-Management-Systeme nicht zu leisten. Daher stellt die Entwicklung eines SLA-fähigen Resource-Management-Systems einen Forschungsschwerpunkt am PC² dar. Ein hier entwickeltes RMS ermöglicht eine Planung der aktuellen und zukünftigen Aufträge sowie eine aktive und flexible Aushandlung der Dienstgüte. Der Benutzer kann definieren, welche Ressourcen benötigt werden, welche Dienstgüte gewährleistet werden muss und wie in einem Fehlerfall reagiert werden soll. Letzteres ist eine besondere Kennzeichnung von kritischen Anwendungen, für die eine hohe Verfügbarkeit gefordert wird. Das RMS kann eventuelle Software- und Hardwareaus-

fälle durch eine Verlagerung der Ausführung auf andere verfügbare Ressourcen innerhalb der gleichen Domäne oder durch Benutzung von weltweit verteilten Ressourcen im Grid kompensieren. Diese Suche kann durchgeführt werden, da dem Provider durch das ausgehandelte SLA das Anforderungsprofil der Anwendung bekannt ist. Es wird folglich eine Anfrage an andere Provider im Grid gestellt, ob sie das vorliegende Anforderungsprofil erfüllen können. Ist diese Suche erfolgreich, so kann das RMS die lokale Berechnung auf die gefundenen Grid-Ressourcen verlagern. Aufgrund der Virtualisierung aller Ressourcen im Grid wird der Benutzer von dieser Verlagerung in der Regel nicht einmal Notiz nehmen. Eine Nutzung von Grid-Ressourcen ist allerdings nur dort sinnvoll, wo die Vereinbarungen des SLAs (z. B. Strafzahlung bei Nicht-Erfüllung) diesen Aufwand rechtfertigen.

Die Verlagerung eines laufenden Prozesses auf andere Ressourcen ist sowohl aus organisatorischer als auch technischer Sicht ein hoch-komplexer Prozess. Zunächst müssen die aktuell laufenden Aufträge ausgewertet und falls möglich – d. h. keine Kollision mit dem gültigen SLA – unterbrochen werden. Die freien Ressourcen werden dann für die Ausführung des höher priorisierten Auftrags gebündelt. Besonders kompliziert wird diese Neuorganisation genau dann, wenn Ergebnisse der laufenden Prozesse als Eingabe für nachfolgende Prozesse benötigt werden. Die technische Realisierung der Auftragsmigration in Grid-Umgebungen ist ein aktuell noch nicht gelöstes Problem und daher Gegenstand der PC² Forschung. Die Fortführung der Berechnung auf einem anderen Knoten erfordert natürlich, dass der Zustand der Berechnung vor dem Ausfall vollständig rekonstruiert werden kann. Üblicherweise wird vorausgesetzt, dass die Berechnung selbständig so genannte „Checkpoints“ erzeugt, die hierfür genutzt werden können. Ist ein Programm hierzu jedoch nicht in der Lage, so ist die gesamte Berechnung bis zum Zeitpunkt des Ausfalls – unter Umständen ein Aufwand von mehreren Tagen oder gar Wochen – verloren. Im Rahmen des EU-geförderten Projekts „HPC4U“, an dem das PC² maßgeblich beteiligt ist, werden daher Mechanismen entwickelt, die das Erzeugen von Checkpoints für beliebige Programme erlauben. Der Anwender muss hierbei keinerlei Änderungen an seinen Applikationen vornehmen, da diese Funktionalität in das RMS integriert wird.

Entsprechend der Anforderungen des Nutzers an die Ausfallsicherheit führt das RMS diese Checkpoints völlig automatisch durch. Dies geschieht in regelmäßigen und ausreichend kurzen Intervallen. Im Fehlerfall – aber auch bei zusammengesetzten Aufträgen – kann somit der aktuelle Status zu einer anderen Ressource in der weltweiten Grid-Umgebung verschoben und noch vor dem Ablauf der Deadline abgearbeitet werden. Bei diesem Migrationsprozess muss das System allerdings sicherstellen, dass die entfernte Grid-Site wirklich kompatibel zu dem aktuell migrierten Auftrag ist. Hierzu ist es erforderlich, dass auf der entfernten Grid-Site Mechanismen zur Verfügung stehen, die im Bedarfsfall die automatische Installation benötigter Bibliotheken oder sonstiger Softwarekomponenten erlauben.

Weitere

Herausforderungen

Innerhalb der praktischen Informatik stellt das Grid-Computing einen hochaktuellen Forschungszweig dar. Weltweit wird an

vielen Problemen in diesem Bereich gearbeitet, wobei die Arbeit durch das *Global Grid Forum (GGF)* [6] koordiniert wird. Neben den geschilderten Problemen müssen noch viele andere Herausforderungen – ein Teil davon ist in Abbildung 2 dargestellt – bewältigt werden. Für einige Probleme wie etwa Sicherheit für die zu übertragenden Daten und Vertrauenswürdigkeit der Grid-Partner existieren Erfahrungen aus der Entwicklung des Webs, die jedoch an die spezifischen Grid-Anforderungen bezüglich Zertifizierung, Authentifizierung und Autorisation angepasst werden müssen. Während bei Webtransaktionen hauptsächlich Daten zu schützen sind, muss hier die Ausführung fremder Programme auf den eigenen Ressourcen geregelt werden.

Durch die geschilderte Entwicklung und Realisierung von SLAs sind alle Erwartungen und Verpflichtungen beider Parteien klar und nachprüfbar festgelegt, so dass die Nutzung des Grid-Computing gemäß den Standards eines kommerziellen Umfelds möglich wird. Der psychologischen Barriere, den eigenen Auftrag einer fremden Umgebung anzuvertrauen, kann allerdings nicht mit technischen Lösungen begegnet werden. Trotz großer Vorteile wie Transparenz und Kosteneinsparungen bilden Grid-Systeme heutzutage noch Insellösungen von „Early Adaptors“. Daher sind gut funktionierende, grid-fähige Beispielanwendungen notwendig, mit denen die erwarteten Zeit- und Kosteneinsparungen nachgewiesen werden können und die dem Grid-Computing zum breiten Durchbruch verhelfen.

Literatur:

- [1] F. J. Corbato and V. A. Vyssotsky, Introduction and overview of the Multics system, AFIPS Conf Proc 27, 185-196, 1965.
- [2] I. Foster: The Anatomy of the Grid: Enabling Scalable Virtual Organizations, Lecture Notes in Computer Science, vol. 2150, 2001.
- [3] Globus Alliance: Globus Toolkit, <http://www.globus.org>.
- [4] UNICORE Forum e.V., <http://www.unicore.org>.
- [5] L.-O. Burchard and M. Hovestadt and O. Kao and A. Keller and B. Linnert, The Virtual Resource Manager: An Architecture for SLA-aware Resource Management, Proc. of the 4th IEEE/ACM International Symposium on Cluster Computing and the Grid, 2004.
- [6] Global Grid Forum, <http://www.ggf.org>.



Dipl.-Inform. Matthias

Hovestadt ist wissenschaftlicher Angestellter im Paderborn Center for Parallel Computing (PC²). In seiner Forschung beschäftigt er sich mit Ressourcenmanagement und Sicherung von Dienstgütern in Grid-Umgebungen.

Einsteigen.



In unserem technisch-orientierten Unternehmen bieten wir laufend interessante Herausforderungen für

**Absolventen^{*)} und Berufserfahrene
der Fachbereiche**

E-Technik, Informatik, Maschinenbau, Mechatronik, Physik, ...

Die dSPACE GmbH ist internationaler Marktführer bei Entwicklungs- und Testwerkzeugen schneller mechatronischer Regelungssysteme wie z.B. ABS oder ESP. Durch die ständige Entwicklung innovativer High-Tech-Produkte wachsen wir seit unserer Gründung 1988 permanent. Deshalb bieten sich immer neue und spannende Aufgaben für unsere Mitarbeiter.

- **Produktmanagement**
- **Hardware-Entwicklung**
- **Software-Entwicklung (GUI, embedded systems)**
- **Anwendungen Echtzeitsimulation**
- **Technische Dokumentation, Marketing und Vertrieb**

Aktuelle Stellenangebote unter www.dspace.de

Bei uns erwarten Sie neueste Technologien, junge, lebendige Projektteams und ein hohes Maß an selbständiger, eigenverantwortlicher Arbeit in einem lockeren, angenehmen Betriebsklima.

*) wir machen keinen Unterschied zwischen Männern und Frauen



dSPACE GmbH · Personalabteilung
Herrn Harald Wilde
Technologiepark 25 · 33100 Paderborn
Tel.: 05251-1638-0 · hwilde@dspace.de



Vermögenssteuer, Verluste und Investitionsentscheidungen

Ökonomische Analyse der Renaissance einer umstrittenen Steuer

Ausgelöst durch die desolate Haushaltslage wird von verschiedenen Seiten die Wiedereinführung einer Vermögenssteuer gefordert. Welche ökonomischen Folgen eine Vermögenssteuer auf das individuelle unternehmerische Investitionsverhalten haben wird, kann durch eine steuerintegrierte investitionstheoretische Analyse herausgearbeitet werden.

Es zeigt sich, dass die Vermögensbesteuerung die Entscheidungen eines Investors erheblich beeinflussen kann. In vielen Fällen verstärkt sie die bestehenden, unerwünschten verzerrenden Wirkungen der Ertragsbesteuerung. Für Investitionsprojekte mit Verlustphasen können zum Teil paradoxe Wirkungszusammenhänge festgestellt werden. Des Weiteren verdeutlichen Überlegungen zur Praktikabilität und zu den Folgen der verfassungsrechtlichen Vorgaben, dass auch hier erhebliche Schwierigkeiten bewältigt werden müssten.



Prof. Dr. rer. pol. Dipl.-Kffr. Caren Sureth ist seit 2004 Inhaberin des Lehrstuhls für Betriebswirtschaftslehre, insbesondere Betriebswirtschaftliche Steuerlehre an der Universität Paderborn. Ihr wissenschaftliches Interesse gilt u. a. der quantitativen Analyse von Steuerwirkungen und dabei vor allem dem Einfluss der Besteuerung auf unternehmerische Kalküle unter Unsicherheit. Unter Anwendung modelltheoretischer Ansätze ergänzt durch Simulationen, Sensitivitätsanalysen sowie durch empirische Arbeiten untersucht sie steuerinduzierte Wirkungszusammenhänge.

Die Vermögenssteuer ist ein Thema, das in regelmäßigen Abständen immer wieder im Mittelpunkt steuerpolitischer Diskussionen steht. Gerade in Zeiten knapper Mittel wird gefordert, dass die Finanzierung des Staates in besonderem Maße von denjenigen zu tragen sei, die über große Vermögen verfügen, um damit gleichzeitig eine Umverteilung von oben nach unten zu bewirken. Dies gilt sowohl für so genannte wohlhabende Privatpersonen als auch für Unternehmungen, die große Kapitalbestände angehäuft haben. Kaum etwas scheint näher zu liegen, als diese Gruppen, die es sich offensichtlich leisten können, durch eine Sondersteuer zu belasten. Schon seit dem Mittelalter ist „soak the rich“ eine Maxime, die auf breite Akzeptanz stößt. Der Gruppe der Befürworter stehen Kritiker einer Vermögenssteuer gegenüber, die neben anderen Aspekten vor allem die Ertragsabhängigkeit dieser Steuer bemängeln. Die jüngsten Zahlen der Steuerschätzung dürften zu einer Verschärfung dieser Auseinandersetzungen beitragen.

Folgen einer Vermögenssteuer

Welche Folgen jedoch tatsächlich durch eine Vermögenssteuer hervorgerufen werden, liegt keineswegs auf der Hand. Unklar ist etwa, ob eine Vermögenssteuer überhaupt die erwarteten und politisch gewünschten Umverteilungseffekte hat. Ungewiss ist zudem, welche ökonomischen Wirkungen und Folgewirkungen auf das Investitionsverhalten von Unternehmungen und damit auf Wachstum und Arbeitsmarkt zu erwarten sind. Strittig ist auch, ob eine Vermögenssteuer überhaupt geeignet ist, die angestrebten Ziele (Umverteilung von Vermögen bei gleichzeitigen positiven Wirkungen auf den Staatshaushalt) zu erreichen bzw. welche Restriktionen hier zu beachten sind.

In Deutschland wurde die Vermögenssteuer zuletzt für das Jahr 1996 erhoben. Das Bundesverfassungsgericht verlangte durch ein

Abb. 1: Formular einer Vermögenssteuererklärung, Württemberg 1764.



Quelle: National Portrait Gallery Pictures Library.

Abb. 2: „John Bull swamped in the flood of new taxes“: – cormorants fishing in the stream, Karikatur von James Gillray, farbiger Druck 1997, Anspielung auf die Steuerpolitik unter dem englischen Schatzkanzler Henry Petty-Fitzmaurice, der den Vermögenssteuersatz von 6,5 auf 10 Prozent erhöhte.

Aufsehen erregendes Urteil im Jahre 1995, dass der Gesetzgeber die unterschiedliche steuerliche Bewertung von Grundvermögen durch die so genannte Einheitsbewertung im Vergleich zu anderem Vermögen beseitigt. Ziel müsse eine verfassungskonforme Vermögensteuer sein. Grundvermögen wurde in vielen Fällen nur mit 10 Prozent des Verkehrswertes in die Berechnung der Vermögensteuer einbezogen, während etwa das Wertpapiervermögen zu 100 Prozent angesetzt wurde. Diese gesetzlich kodifizierten Bewertungsvorschriften verstoßen gegen den Gleichheitsgrundsatz des Grundgesetzes. Der Aufforderung des Verfassungsgerichtes ist der Gesetzgeber nicht nachgekommen. Aus diesem und weiteren Gründen ist die Erhebung der Vermögensteuer ab dem 1.1.1997 ausgesetzt worden. Obwohl Nachbesserungen in der Bewertung von Vermögen stattgefunden haben, hat der Bundesfinanzhof im Zusammenhang mit der Erbschaftsbesteuerung die Verfassungsmäßigkeit der Bewertungsvorschriften dem Bundesverfassungsgericht zur Klärung vorgelegt. Dies ist ein Indiz dafür, dass die Einwände der Verfassungsrichter bisher nicht beseitigt worden sind [4].

Die kleine Steuer

Die Erhebung und Entrichtung der Vermögensteuer rief in der Vergangenheit hohe Kosten hervor, die auf etwa ein Drittel der Einnahmen aus der Vermögensteuer geschätzt wurden. Insgesamt betrug das Vermögensteueraufkommen ohne Berücksichtigung dieser Kosten im Jahr 1996 mit umgerechnet 4,5 Mrd. Euro lediglich ca. 1 Prozent des Gesamtsteueraufkommens dieses Jahres. Damit wurde die Vermögensteuer zu Recht als kleine Steuer bezeichnet.

Obwohl die Vermögensteuer ein besonders ungünstiges Verhält-

nis von Aufkommen zu Vollzugskosten aufweist und erhebliche Probleme hinsichtlich der realitätsnahen Bewertung bestimmter Arten von Vermögen birgt, erfuhr die kontroverse Diskussion um die Wiederbelebung dieser Steuer in Deutschland im Umfeld des Bundestagswahlkampfes im Jahre 2002 eine Renaissance. Ausgelöst durch die desolante Haushaltslage wurde nicht nur die Vermögensteuer erneut Gegenstand der öffentlichen Diskussion, sondern u. a. auch der Ruf nach einer Beschränkung der steuerlichen Verlustverrechnungsmöglichkeiten, insbesondere von Kapitalgesellschaften. Eine Neuregelung der Beschränkung der Verlustverrechnung, die so genannte Mindestbesteuerung, ist bereits zum 1.1.2004 in Kraft getreten ist. Die Debatte um eine Vermögensteuer ist hingegen noch nicht beendet. Die Beschlüsse auf den Parteitagen der Regierungsparteien im Sommer 2003



Abb. 3: Kaiser Karls VI. Ankündigung einer neuen Vermögenssteuer, 1713.

Quelle: Finanzgeschichtliche Sammlung der Bundesfinanzakademie.

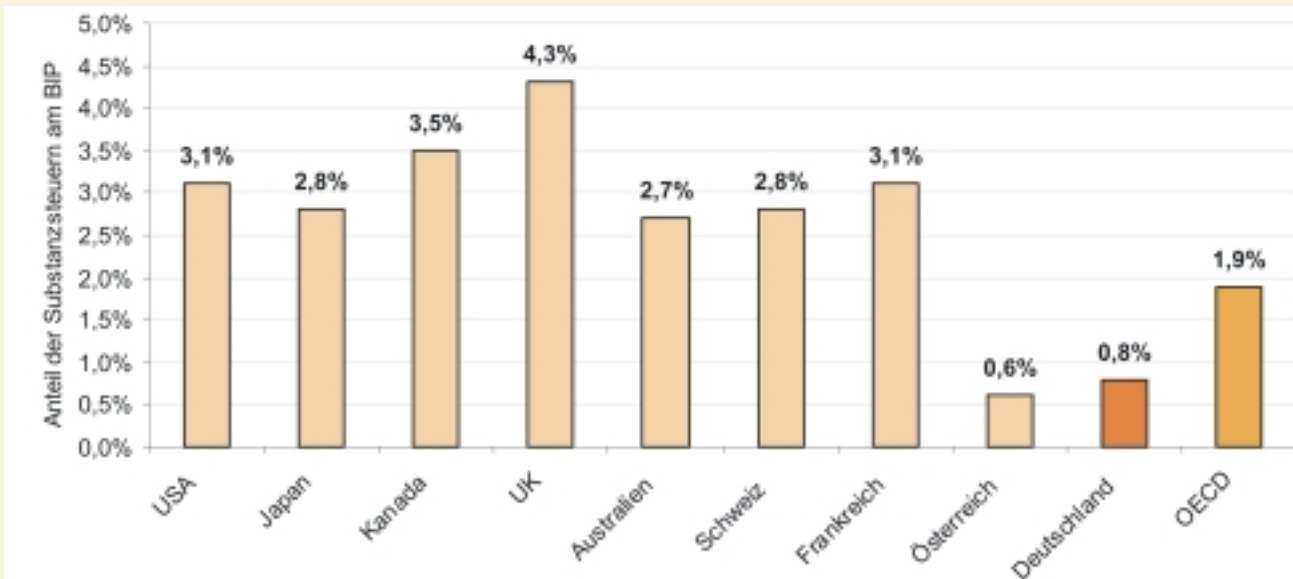


Abb. 4: Substanzsteueranteil am BIP in ausgewählten OECD-Staaten im Jahr 2001.

haben verdeutlicht, dass es sich bei der Diskussion über die Wiederbelebung der Vermögensteuer nicht um ein rein akademisches Thema handelt. Während man mittlerweile in der SPD weitgehend Abstand von diesem Vorschlag genommen hat, befürwortete eine Gruppe der Bundesdelegiertenkonferenz von Bündnis 90/Die Grünen im Oktober 2004 nach wie vor die Wiedereinführung einer (modifizierten) Vermögensteuer.

Der folgende Beitrag konzentriert sich im Wesentlichen auf die Folgen einer Vermögensbesteuerung für unternehmerische Investitionsentscheidungen. Anhand einfacher Szenarien werden mögliche Wirkungen einer Vermögensteuer an sich und in Kombination mit einer Beschränkung der steuerlichen Geltendmachung von Verlusten dargestellt und erläutert. Abgesehen von der aktuellen Brisanz dieses Themas, liefern diese Ergebnisse grundsätzliche Einblicke in die Wirkungsweise steuerrechtlicher Regelungen. Damit handelt es sich um ein Thema, das nicht nur tagespolitisch, sondern vor allem auch wissenschaftlich im Hinblick auf eine fundierte ökonomische Analyse steuerrechtlicher Regelungen interessant ist.

Steuerrechtliche Rahmenbedingungen

Wirft man zur Orientierung einen Blick auf die Statistik der OECD (Abbildung 4), stellt man fest, dass Deutschland im internationalen Vergleich mit 0,8 Prozent Substanzsteueranteil am Bruttoinlandsprodukt eine relativ geringe Quote aufweist. Darin enthalten sind neben einer etwaigen Vermögensteuer, auch Erbschaft- und Schenkungsteuern sowie die Liegenschaftsteuern, wie z. B. Grundsteuern, der jeweiligen Staaten. Diese Statistik wird gerne als Argument für eine Wiederbelebung der Vermögensteuer angeführt. Hierbei werden jedoch steuerliche Unterschiede in den Ertragsteuerbelastungen der Länder sowie die Investitionswirkungen einer Vermögensteuer ignoriert. International erheben nur noch wenige OECD-Staaten eine Vermögensteuer auf das Vermögen natürlicher Personen. In Europa sind dies nur noch Finnland, Frankreich, Island, Luxemburg, Norwegen, Spanien, Schweden und die Schweiz, davon erheben Island, Luxemburg und die Schweiz auch eine Vermögensteuer auf das Vermögen juristischer Personen.

Zuletzt wurde in Deutschland eine Vermögensteuer in Höhe von 0,5 Prozent auf das produktive Vermögen bzw. 1 Prozent auf das

nichtproduktive Vermögen natürlicher Personen erhoben sowie 0,6 Prozent auf das Betriebsvermögen von Kapitalgesellschaften. Investitionsentscheidungen werden durch Verluste und deren steuerliche Behandlung stark beeinflusst. Da in Phasen konjunktureller Schwäche Verlustsituationen relativ häufig auftreten, ist es sinnvoll, eine Analyse vermögenssteuerlicher Effekte um verlustinduzierte ertrag- und substanzsteuerliche Wirkungen zu ergänzen. Dies gilt vor allem deshalb, weil Verluste das Vermögen eines Steuerpflichtigen reduzieren und damit neben ertragsteuerlichen Wirkungen unmittelbar einen Einfluss auf die vermögenssteuerliche Bemessungsgrundlage haben. Gegenwärtig gestattet es das deutsche Steuerrecht, Verluste eine Periode zurückzutragen und mit Gewinnen der Vorperiode zu verrechnen. Verbleibende Verluste können in zukünftige Perioden vorgetragen werden. Seit dem 1.1.2004 ist die Verrechnung von Verlusten aus Vorjahren mit laufenden Einkünften bis zu einem Gesamtbetrag der Einkünfte von 1.000.000 Euro unbeschränkt möglich, darüber hinaus jedoch auf 60 Prozent des übersteigenden Gesamtbetrags der Einkünfte begrenzt (Mindestbesteuerung). Diese Maßnahme soll eine Verstetigung der Staatseinnahmen bewirken und insbesondere die extremen Einbrüche des Körperschaftssteueraufkommens eindämmen (Körperschaftssteueraufkommen 2000: 23.190 Mio. Euro, 2001: -426 Mio. Euro, 2002: 2.864 Mio. Euro). Die Mindestbesteuerung verschärft die bereits existierende asymmetrische steuerliche Behandlung von Gewinnen und Verlusten und benachteiligt Verlustunternehmen damit in besonderem Maße. Bei Investitionsprojekten mit im Zeitablauf negativen und positiven Jahresergebnissen im Vergleich zu Investitionen mit gleichmäßigen (niedrigeren) positiven Ergebnissen, aber identischer Rendite vor Berücksichtigung von Steuern (Vorsteuerrendite), erfahren die zeitweise Verluste erwirtschaftenden Projekte einen Nachteil. Im günstigsten Fall kommt es durch die Verlustverrechnungsbeschränkung lediglich zu einem Zinsnachteil für Investitionen mit Verlustperioden [1], [8].

Die folgende Analyse konzentriert sich zunächst auf die Wirkungen einer Vermögensteuer auf unternehmerische Investitionsentscheidungen in Gewinnsituationen [3]. Dieses Szenario wird anschließend um Investitionsprojekte mit Verlustphasen erweitert, bei denen die Regelungen der Mindestbesteuerung wirksam werden. Auf diese Weise werden die Wechselwirkungen von Vermögensbesteuerung, Verlusten und Verlustverrechnungsbe-

schränkungen herausgearbeitet. Dieses Vorgehen (Partialanalyse) erlaubt es zwar nicht, gesamtwirtschaftliche Konsequenzen abzuschätzen, jedoch können spezielle einzelwirtschaftliche Wirkungszusammenhänge transparent gemacht werden und damit Aussagen über tendenzielle gesamtwirtschaftliche Folgen gewonnen werden. Die Untersuchung kann damit zur Versachlichung der öffentlichen Diskussion beitragen.

Steuerökonomische Analyse

Üblicherweise ist es gewünscht, dass ein Investor in solche Projekte investiert, die eine möglichst hohe Rendite versprechen und auf diese Weise größtmögliche Wachstumseffekte hervorruhen. Steuern dürfen – wenn durch den Staat nicht explizit Lenkungsabsichten verfolgt werden – nicht dazu führen, dass Ressourcen (Investitionskapital) in Projekte fließen, die ohne eine steuerliche Privilegierung unwirtschaftlich wären. Da das Ziel einer Vermögensteuer grundsätzlich nicht die Lenkung von Investitionskapital sein sollte, ist eine Vermögensteuer so auszugestalten, dass die Entscheidungen von Investoren nicht bzw. kaum beeinflusst werden. Um den Einfluss einer Vermögensteuer zu isolieren, bieten einfache Szenarien große Vorteile. Sie sind intuitiv leicht nachvollziehbar, bergen jedoch den Nachteil, dass die erzielbaren Ergebnisse nicht verallgemeinerbar sind, da eine Vielzahl von Einflussfaktoren vernachlässigt wird. Lediglich Tendenzaussagen sind möglich.

Als Szenario wird im Folgenden ein Investor betrachtet, der sein Kapital entweder in einem Investitionsprojekt in einer Unternehmung (Realinvestition) anlegt oder aber diese Investition unterlassen kann. Im Fall der Unterlassungsalternative verdient der Investor Zinsen aus der Anlage seiner Mittel in festverzinsliche Kapitalmarktpapiere. Beide Alternativen erwirtschaften annahmegemäß eine übereinstimmende Vorsteuerrendite von 10 Prozent. Ausgehend von einem vollkommenen Kapitalmarkt unter Sicherheit werden simultan Ertragsteuern, wie etwa eine Einkommensteuer, Körperschaftsteuer oder Gewerbesteuer, sowie eine Vermögensteuer in den Kalkül integriert. Im Vergleich zu den Wirkungen einer Vermögensteuer auf die Kapitalmarktanlage zeigt die Integration einer Vermögensteuer folgende Wirkungen auf die erzielbare Nachsteuerrendite:

In Abhängigkeit von der Nutzungsdauer eines abnutzbaren und damit abschreibbaren Investitionsgutes zeigt sich für den Fall zeitkonstanter Einzahlungsüberschüsse bei Anwendung der derzeit gesetzlich zulässigen Abschreibungsverfahren, dass eine Vermögensteuer zu einer relativen Begünstigung der Realinvestition führt. Eine Vermögensteuer mindert die Rendite eines Investitionsvorhabens zwar absolut, im Vergleich zur Kapitalmarktanlage genießt die Realinvestition jedoch relative Vorteile (vgl. Abbildung 5). Die Grafik veranschaulicht, dass die relative Renditeänderung in Folge der Ertrags- und Vermögensbesteuerung größer als Null ist. Die Rendite nach Steuern erhöht sich im Fall der Realinvestition relativ zur Alternativanlage. Das ist bei allen betrachteten Nutzungsdauern der Fall.

Diese investitionsfördernde Wirkung einer Vermögensteuer ist ursächlich auf eine Subventionierung durch die bestehenden Abschreibungsvorschriften zurückzuführen. Sie bewirken unter Berücksichtigung der zeitlichen Struktur der Zahlungen eine geringere steuerliche Belastung der Zahlungsüberschüsse aus der Realinvestition als die Besteuerung der Zinsen aus der Kapital-

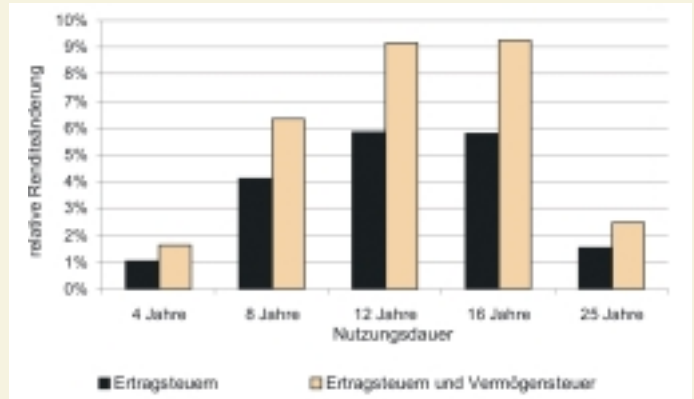


Abb. 5: Relative Rentabilitätswirkungen bei im Zeitablauf konstanten Einzahlungsüberschüssen.

marktanlage. Dieser als Steuerparadoxon bekannte Abschreibungseffekt wirkt bereits bei den Ertragsteuern, wie man in Abbildung 6 anhand der dunklen Säulen erkennen kann [5], und wird durch die Vermögensbesteuerung verstärkt. Den negativen Liquiditätseffekten einer zusätzlichen Besteuerung steht in diesem Fall somit eine relative Renditeverbesserung gegenüber.

Es lässt sich jedoch zeigen, dass dieser Effekt vom unterstellten zeitlichen Verlauf der periodischen Überschüsse aus dem Investitionsprojekt abhängt und bei einer zusätzlich erhobenen Vermögensteuer besonders sensibel auf Veränderung der zeitlichen Struktur reagiert. Während wachsende Zahlungsüberschüsse das Ausmaß der Subventionierung erhöhen, kommt es bei fallenden Verläufen (z. B. 15 Prozent negatives Wachstum bei gleich bleibender Vorsteuerrendite) zur relativen Benachteiligung der Realinvestition. Fallende Zahlungsüberschüsse implizieren hohe Gewinne zu einem frühen Zeitpunkt und damit frühe hohe Steuerzahlungen. Dies bewirkt letztlich einen relativen Zinsnachteil für die Realinvestition. Da die Abschreibungen nicht nur die ertragsteuerliche Bemessungsgrundlage reduzieren, sondern auch das Vermögen und damit die vermögenssteuerliche Bemessungsgrundlage beeinflussen, wird dieser Effekt durch die Einführung einer Vermögensteuer verstärkt.

Eindeutige Aussagen darüber, ob eine Vermögensteuer unter Rentabilitätsaspekten Investitionen fördert oder behindert oder sogar neutral behandelt, sind nicht möglich. Dieses Untersuchungsszenario verdeutlicht bereits, dass die Ergebnisse besonders sensitiv hinsichtlich des zeitlichen Verlaufs der Zahlungsüberschüsse und Abschreibungen sind. Die unterstellten Zahlungsüberschüsse stellen jedoch nur Beispiele dar, welche diese Steuerwirkungen besonders deutlich machen. Üblicherwei-

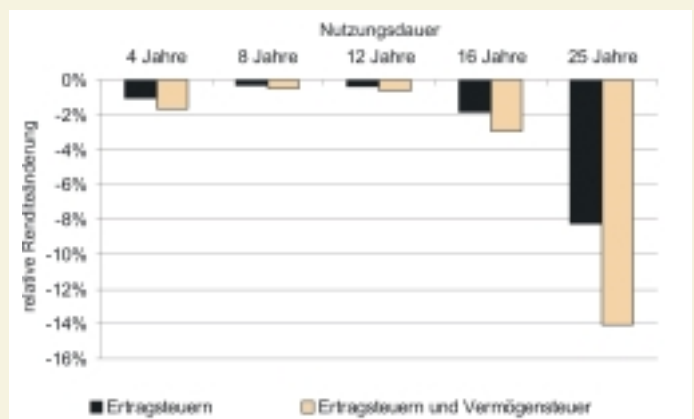


Abb. 6: Relative Rentabilitätswirkungen der Besteuerung bei im Zeitablauf fallenden Einzahlungsüberschüssen.

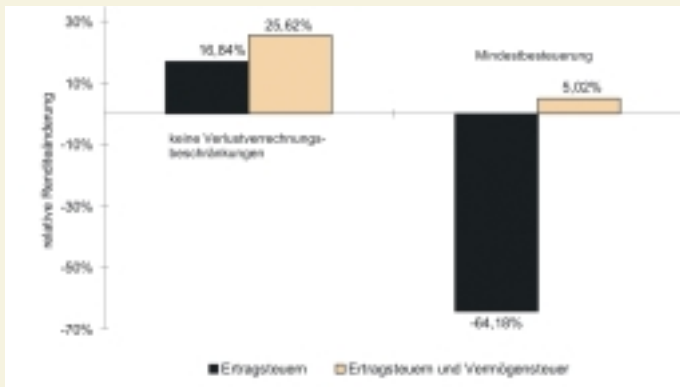


Abb. 7: Renditewirkungen der Mindestbesteuerung.

se verlaufen die Einzahlungsüberschüsse aus einem Investitionsprojekt im Zeitablauf nicht homogen, sondern werden durch verschiedene Faktoren beeinflusst. Sie weisen in der Regel sowohl Phasen mit zunehmenden als auch fallenden Überschüssen auf.

Möchte man den Einfluss steuerlicher Verlustverrechnungsbeschränkungen mit in die Untersuchung einbeziehen, liegt es nahe, wiederum eine Investition mit gleicher Vorsteuerrendite wie die Kapitalmarktanlage zu betrachten, die beispielsweise zu Beginn des Investitionszeitraums zunächst steuerliche Verluste und später entsprechend hohe Überschüsse aufweist. Es zeigt sich, dass die phasenverschobene steuerliche Berücksichtigung von Verlusten im Einzelfall den Vermögenssteuereffekt auf die Investitionsentscheidung des Investors kompensieren kann. Die Abbildung 7 veranschaulicht, dass im Einzelfall durch die Vermögensbesteuerung aus einer nach Ertragsteuern zunächst unvorteilhaften Investition, die ein Investor nicht durchführen würde, sogar eine vorteilhafte Investition resultieren kann.

Unterstellt man, dass alle Verluste bis zum Ende der Nutzungsdauer mit Gewinnen verrechnet werden können, kommt es durch die zeitliche Streckung des Verlustverrechnungspotenzials zu einem Zinsnachteil. Die Vermögensbesteuerung kann diesen Nachteil zum Teil auffangen, da der Verlust wiederum die Bemessungsgrundlage der Vermögenssteuer reduziert. Unter diesen Voraussetzungen kann sich der Vorteil aus den Abschreibungen bei wachsenden Zahlungsüberschüssen in besonderem Maße entfalten und bewirkt einen besonders großen relativen steuerlichen Vorteil im Vergleich zur reinen Ertragsbesteuerung. Diese Beispielrechnungen dokumentieren exemplarisch, dass die Vermögensbesteuerung nicht zwangsläufig dazu führt, dass die Steuerpflichtigen hieraus unter Rentabilitäts Gesichtspunkten besondere Nachteile erfahren. In vielen Fällen verstärkt sie (unerwünschte) Entscheidungswirkungen, die bereits im Fall einer reinen Ertragsbesteuerung beobachtet werden können. Man darf allerdings bei der hier gewählten relativen Betrachtungsweise nicht vergessen, dass es selbstverständlich zu absoluten Renditeeinbußen für den Investor kommt. Dies gilt für die Real- wie die Kapitalmarktanlage. Auch wenn die unterstellten Szenarien keine Verallgemeinerungen erlauben, so zeigen sie dennoch, dass eine Vermögenssteuer bereits unter idealisierten Bedingungen sehr heterogen wirken kann. Sie wirkt für bestimmte Investitionsprojekte investitionsfördernd und zugleich für andere investitions-hemmend. Die Wirkungsrichtung hängt maßgeblich von der zeitlichen Struktur der Zahlungen sowie der steuerlichen Bemessungsgrundlagen und damit vom einzelnen Investitionsobjekt ab.

Eindeutig ist hingegen, dass sich eine Vermögenssteuer mindernd auf die Liquidität auswirkt, da in der Realität die Bedingungen des vollkommenen Kapitalmarktes und damit unbeschränkte Kapitalaufnahmemöglichkeiten nicht bestehen. Dieser Aspekt ist besonders in Verlustphasen von erheblicher Bedeutung.

Verfassungsrechtliche Anforderungen

Möchte man ungeachtet dieser Effekte und ungeachtet der Tatsache, dass die Erhebung einer Vermögenssteuer hohe Kosten und bislang unlösbare Bewertungsprobleme birgt, aus politökonomischen Erwägungen heraus dennoch eine Vermögenssteuer einführen, so muss man sich nicht zuletzt mit den verfassungsrechtlichen Restriktionen auseinander setzen. Das Urteil des Bundesverfassungsgerichtes im Jahr 1995 fordert den Gesetzgeber auf

1. eine realitätsgerechtere Bewertung zu implementieren,
2. dafür zu sorgen, dass die Vermögenssteuer keinen Eingriff in die Vermögenssubstanz und keine Belastung des Existenzminimums der Steuerpflichtigen hervorruft und
3. dass die steuerliche Gesamtbelastung der Einkünfte durch die Vermögensbesteuerung nicht zu mehr als einer hälftigen Teilung der Erträge zwischen Fiskus und Steuerpflichtigen führt (Halbteilungsgrundsatz).

Um die Wirkungen des Halbteilungsgrundsatzes zu veranschaulichen, ist die Betrachtung einerseits von Steuerpflichtigen mit niedrigem Vermögen und andererseits von Steuerpflichtigen mit hohem Vermögen hilfreich. Unterstellt man einen Vermögenssteuersatz von 1 Prozent und konzentriert sich auf das Verhältnis von Steuerbelastung aus Ertrag- und Vermögenssteuern zu Vorsteuerrendite, wird die Bedeutung dieses Grundsatzes transparent.

In Abbildung 8 kann man erkennen, dass die Vermögensbesteuerung bei Spitzenverdienern dem Halbteilungsgrundsatz regelmäßig nicht genügen kann. Bezieht man Solidaritätszuschlag und Kirchensteuer mit ein, werden auch hochrentable Investitionen mit mehr als 50 Prozent steuerlich belastet. Die Erhebung einer Vermögenssteuer für Vermögen mit geringen Erträgen verstößt ebenfalls regelmäßig gegen den Halbteilungsgrundsatz, da hier das Verhältnis der Vermögenssteuer auf das Vermögen zu den geringen Erträgen insgesamt zu einer besonders hohen Gesamtsteuerbelastung der Erträge führt. Untersuchungen weiterer Vermögens- und Ertragskombinationen bestätigen schließlich, dass verfassungsrechtlich nur eine Vermögensbesteuerung für ertragreiche kleine und mittlere Vermögen bei gleichzeitig niedri-

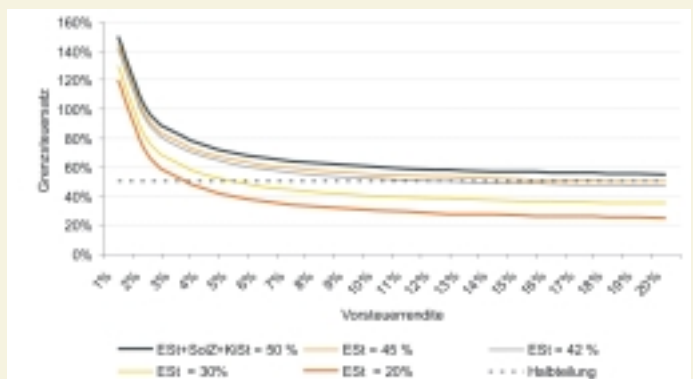


Abb. 8: Zusammenhang von Rendite, Ertrags- und Substanzsteuerbelastung und Halbteilungsgrundsatz.

gen oder mittleren Einkünften verfassungskonform sein kann. Berücksichtigt man die zweite Forderung des Bundesverfassungsgerichtes, muss der Kreis der möglichen Steuerpflichtigen abermals verringert werden. Die resultierenden möglichen Besteuerungsfälle stimmen somit nicht mit den in der politischen Diskussion angeführten gewünschten Belasteten überein.

Fazit und Ausblick

Es hat sich gezeigt, dass die traditionelle Vermögensbesteuerung unter Umständen einen erheblichen Einfluss auf die Entscheidungen eines Investors ausüben kann. In vielen Fällen verstärkt sie die an sich unerwünschten verzerrenden Wirkungen der Ertragsbesteuerung. Im Einzelfall lässt sich zeigen, dass eine Vermögensteuer „zufällig“ zu derselben Investitionsentscheidung wie im Vorsteuerergebnis führen kann und damit neutral wirkt. Diese Neutralität ist das Ergebnis kompensierender Effekte aus abschreibungsbedingten Steuerersparnissen und Zinseffekten, die gegebenenfalls durch die komplexen, zum Teil paradoxen Wirkungen von Verlustverrechnungsbeschränkungen zusätzlich beeinflusst werden. Welche gesamtwirtschaftlichen Effekte auftreten, lässt sich auf der Grundlage dieses Partialmodells nicht ableiten. Festgehalten werden kann jedoch, dass bereits unter einfachen Modellbedingungen erhebliche Entscheidungswirkungen in unterschiedliche Richtungen denkbar sind. Tendenziell kann auf Grundlage der hier angefügten Berechnungen gefolgert werden, dass erhebliche Effizienzwirkungen von einer Vermögensbesteuerung ausgehen dürften.

Will man Untersuchungen in einer Modellwelt unter Unsicherheit durchführen, um damit größere Realitätsnähe des Analyse-rahmens zu erreichen, zeigt sich, dass die Modelle einen erheblich höheren Komplexitätsgrad aufweisen [6], [7], [8]. Die Anzahl der sensitiven Parameter steigt an und in vielen Fällen sind nur noch bedingte Ergebnisse ableitbar. Diese lassen sich oft nur schwer intuitiv erschließen. Ihr Beitrag für den politischen Willensbildungsprozess darf daher nicht überschätzt werden.

Die Überlegungen zur Praktikabilität und zu den Folgen der verfassungsrechtlichen Vorgaben machen deutlich, dass von dieser Seite erhebliche Schwierigkeiten bewältigt werden müssen, für die bislang keine überzeugenden Lösungsansätze vorliegen. Um den Kreis derjenigen Steuerpflichtigen, die verfassungskonform mit Vermögensteuer belastet werden dürfen, zu erweitern, könnte man die Möglichkeit der Anrechnung der Vermögensteuer auf die Körperschaft- und/oder Einkommensteuer vorsehen, was jüngst diskutiert wird. Ohne eine Beschränkung der Anrechnung der Vermögensteuer würden jedoch sowohl die Aufkommenswirkungen als auch etwaige angestrebte Umverteilungswirkungen verloren gehen [2]. Lösungsansätze, die zugleich der Verfassung und den erwähnten Praktikabilitätsproblemen gerecht werden, sind bisher nicht erkennbar.

Glossar

Grenzsteuersatz:

Der Grenzsteuersatz gibt die steuerliche Belastung an, die auf eine zusätzliche, infinitesimal kleine Einheit der steuerlichen Bemessungsgrundlage entfällt.

Halbteilungsgrundsatz:

Der Halbteilungsgrundsatz besagt, dass die steuerliche Gesamtbelastung der Einkünfte nicht zu mehr als etwa einer hälftigen

Teilung zwischen Fiskus und Steuerpflichtigen führen darf.

Partialanalyse

Untersuchung auf einem Teilmarkt unter Vernachlässigung aller Wirkungen auf anderen (Teil-)Märkten.

Produktives Vermögen:

Produktives Vermögen natürlicher Personen umfasst land- und forstwirtschaftliches Vermögen, Betriebsvermögen von Personengesellschaften oder Einzelunternehmungen sowie Beteiligungen an Kapitalgesellschaften.

Rendite:

Die Vorsteuer-Rendite r (Baldwin-Rendite oder so genannte Vermögensrentabilität) gibt die durchschnittliche geometrische Verzinsung des Anfangsvermögens (eingesetztes Kapital) bezogen auf den Planungshorizont T an:

$$r = \sqrt[T]{\frac{\text{Endvermögen vor Steuern}}{\text{Anfangsvermögen}}} - 1. \text{ Für die Nachsteuerrendite folgt}$$

$$\text{entsprechend: } r^s = \sqrt[T]{\frac{\text{Endvermögen nach Steuern}}{\text{Anfangsvermögen}}} - 1$$

Verzerrungen:

Steuerinduzierte Verzerrungen liegen vor, wenn Steuerpflichtige versuchen, durch Verhaltensänderungen ihre Steuerlast zu verringern.

Vollkommener Kapitalmarkt unter Sicherheit:

Das Modell des vollkommenen Kapitalmarktes ist gekennzeichnet durch symmetrisch verteilte Informationen, unendlich schnelle Reaktionsgeschwindigkeit, keine Transaktionskosten und unbeschränkte Kapitalaufnahme- und -anlagemöglichkeiten, vollkommene Konkurrenz für alle Marktteilnehmer, sichere Umweltzustände.

Literatur

- [1] Auerbach, Alan J.; Poterba, James M. (1987): Tax-Loss Carryforwards and Corporate Tax Incentives, in: Feldstein, Martin (ed.), The Effects of Taxation on Capital Accumulation, Chicago, London 1987, 305-342.
- [2] Bach, Stefan; Haan, Peter; Maiterth, Ralf; Sureth, Caren (2004): Modelle für die Vermögensbesteuerung von natürlichen Personen und Kapitalgesellschaften – Konzepte, Aufkommen, wirtschaftliche Wirkungen, in: DIW Berlin: Politikberatung kompakt 2004, Berlin.
- [3] Georgi, Andreas A. (1986): Steuern in der Investitionsplanung. Eine Analyse der Entscheidungsrelevanz von Ertrag- und Substanzsteuern, Hamburg 1986.
- [4] Haegert, Lutz; Maiterth, Ralf (2002): Zum Ausmaß der steuerlichen Unterbewertung von Grundstücken nach geltendem Recht und bei Anwendung der Reformvorschriften eines Gesetzentwurfs von fünf Bundesländern, in: Steuer und Wirtschaft, 2002, 79. Jg., 248-260.
- [5] König, Rolf; Sureth, Caren (2002): Die ökonomischen Auswirkungen der Änderungen der steuerlichen Abschreibungsmodalitäten, in: Die Betriebswirtschaft, 2002, 62. Jg., 260-272.
- [6] Niemann, Rainer; Sureth, Caren (2004): Tax Neutrality under Irreversibility and Risk Aversion, in: Economics Letters, 2004, vol. 84, 43-47.
- [7] Sureth, Caren (2002): Partially irreversible investment decisions and taxation under uncertainty – a real option approach, in: German Economic Review, 2002, vol. 3, 185-221.
- [8] Wijnbergen, Sweder van; Estache, Antonio (1999): Evaluating the minimum asset tax on corporations: an option pricing approach, in: Journal of Public Economics, 1999, vol. 71, 75-96.

Schnelle Mikrochips

Preiswerter Anschluss an die Datenautobahn

Seit Jack Kilby 1958 bei der Firma Texas Instruments den ersten integrierten Schaltkreis entwickelte – er erhielt hierfür im Jahr 2000 den Nobelpreis für Physik – haben die so genannten Mikrochips nahezu alle Lebensbereiche des Menschen nachhaltig verändert. Mikroprozessoren begegnen uns selbst in Haushaltsgeräten. Personalcomputer bieten eine Rechenleistung, die noch vor wenigen Jahren nur ausgewählten Forschungseinrichtungen zur Verfügung stand. Nunmehr erleben wir die Revolutionierung der Kommunikationstechnik und deren Verschmelzung mit dem Computer, der sich vom Rechner zum mobilen multimedialen Zentrum jedes Haushalts entwickelt.

Datenübertragung im Zeitmultiplex

Die Rückgratverbindungen des weltumspannenden Kommunikationsnetzes werden heute durch Glasfaserverbindungen gebildet. Dabei können über eine einzige Faser ca. 100 optische Kanäle übertragen werden, die sich in der Wellenlänge des für die Übertragung genutzten Infrarotlichtes unterscheiden. In jedem dieser Kanäle können wiederum bis zu 40 Milliarden Bits pro Sekunde übertragen werden, was bereits der Datenrate von etwa 3 Millionen Telefongesprächen oder 20 000 Videokanälen entspricht. Um diese enorme Übertragungskapazität auslasten zu können, müssen viele, relativ langsame elektronische Datenkanäle zunächst im Zeitmultiplex zu einem schnellen Kanal zusammengefasst werden. Diese Schnittstelle der Elektronik zur optischen Übertragungsstrecke stellt daher ein aktuelles Arbeitsgebiet der Hochfrequenzelektronik dar [1]. Das Prinzip ist in Abbildung 1 dargestellt.

Auf der Senderseite schaltet zunächst ein so genannter Multiplexer (MUX) die an n verschiedenen Eingängen einlaufenden Datenbits nacheinander auf den Ausgang, so dass hier die Datenrate gegenüber den Eingängen um gerade den Faktor n vergrößert ist. Dieser Datenstrom wird nun genutzt, um mit Hilfe eines Modulators das Licht eines im Dauerbetrieb arbeitenden Lasers an- bzw. abzuschalten. Da die heute verwendeten Modulatoren relativ hohe Steuerspannungen erfordern, wird das elektrische Signal zuvor in einem Modulatortreiber (MD)

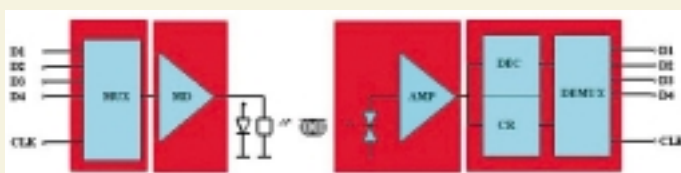


Abb. 1: Optoelektronische Übertragungsstrecke im elektrischen Zeitmultiplex (ETDM).



Prof. Dr.-Ing. Andreas Thiede gründete mit seiner Berufung an die Universität Paderborn im Jahre 1999 das Fachgebiet Höchstfrequenzelektronik, das er seither leitet. Seine Forschungsschwerpunkte sind applikationsspezifische integrierte Halbleiterschaltkreise (ASIC) mit gemischt analogen und digitalen Funktionen für hochfrequenztechnische und insbesondere optoelektronische Anwendungen.

verstärkt. Die so erzeugten kurzen Lichtimpulse werden nun über eine Glasfaser übertragen, im Empfänger zunächst mit Hilfe einer Photodiode wieder in elektrische Impulse umgewandelt und verstärkt. Anschließend muss aus dem Datenstrom das Taktsignal (CLK) zurückgewonnen werden (CR). Mit Hilfe des Taktsignals kann dann zum jeweils richtigen Zeitpunkt eine Entscheidung (DEC) gefällt werden, ob es sich bei dem jeweils empfangenen Bit um eine 1 oder 0 gehandelt hat, und der digitale Datenstrom schließlich mit einem Demultiplexer auf wiederum n elektronische Kanäle mit n -fach geringerer Datenrate verteilt werden. Bislang werden die für ein derartiges System erforderlichen Komponenten in sehr aufwendigen und daher teuren Technologien auf der Grundlage von Verbundhalbleitern wie GaAs und InP hergestellt. Mit dem Einzug der zunächst für wenige Weitverkehrsstrecken eingesetzten optischen Übertragungstechnik in den innerstädtischen Bereich ist jedoch zwingend eine deutliche Reduzierung der Kosten verbunden. Diese kann nur durch den Einsatz der Standard CMOS-Technologie auf Si-Basis, der Grundlage von weit über 90 Prozent aller weltweit hergestellten integrierten Schaltkreise, erreicht werden. Die CMOS-Technologie, in den 70er Jahren ursprünglich für leistungsarme digitale Anwendungen wie Armbanduhren und Taschenrechner entwickelt, konnte sich durch eine kontinuierliche Verringerung der typischen Dimensionen von mehreren Mikrometern bis zu heute weniger als $0.1\ \mu\text{m}$ nunmehr auch für einen Einsatz in der Hochfrequenztechnik qualifizieren. Diese Entwicklung kann von jedermann leicht anhand der Taktfrequenzen und der Größe der Hauptspeicher handelsüblicher Personalcomputer nachvollzogen werden. Doch die Si CMOS-Technologie ist und bleibt eine vor allem für digitale Schaltkreise, wie etwa Mikroprozessoren und Speicher, optimierte Technologie, so dass der Entwickler analoger Schaltkreise schwerwiegende Hindernisse zu überwinden hat. Ein Problem für den Schaltungsentwurf stellt die mit der Minia-

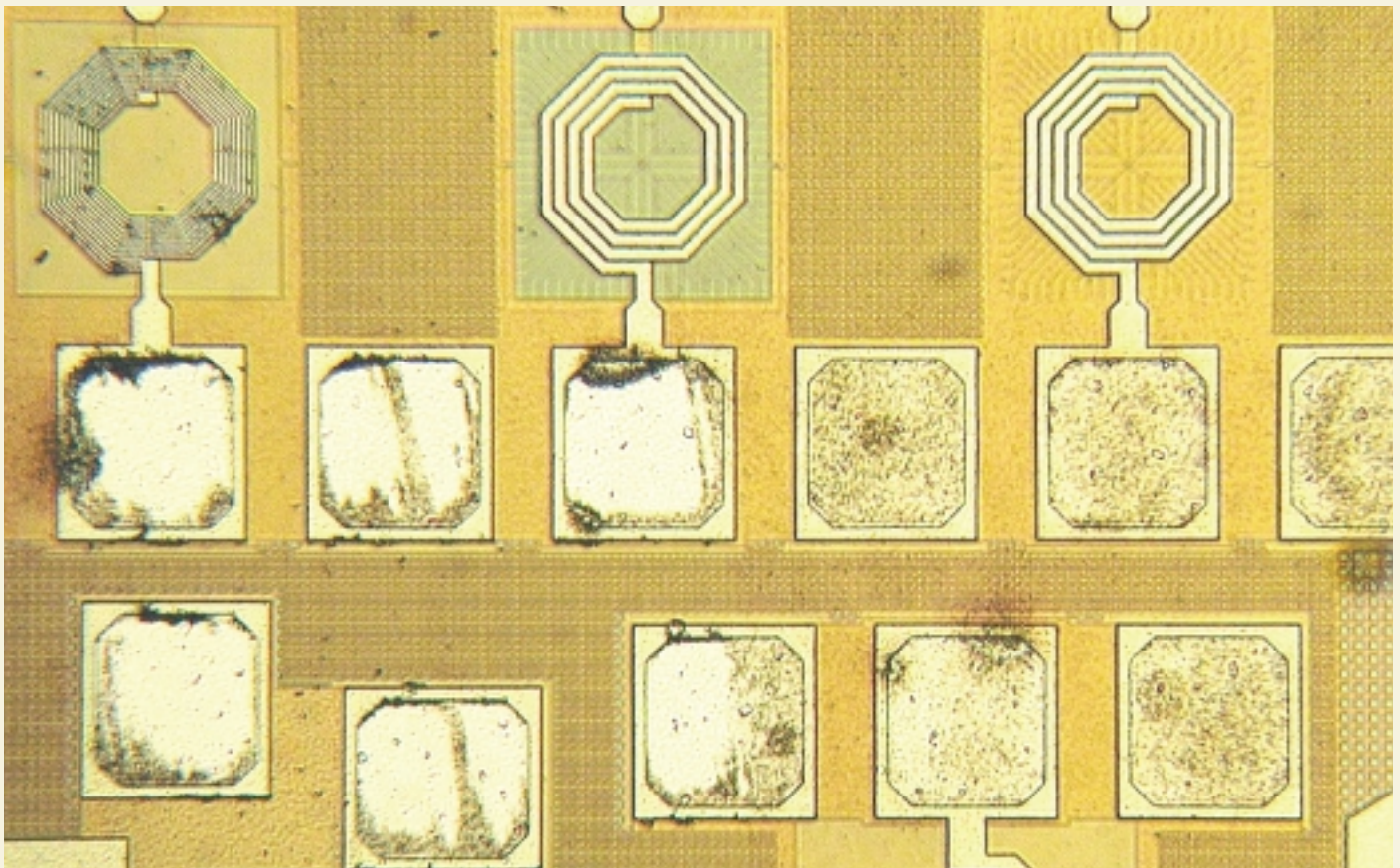


Abb. 2: Integrierte Spiralspulen.

turisierung einhergehende Absenkung der Durchbruchspannungen und damit der maximal zulässigen Betriebsspannung dar, so dass Standardtechniken nicht mehr einsetzbar sind und neue Lösungen gesucht werden müssen. Ein weiteres Problem stellen die hochdotierten und damit sehr gut leitfähigen Substrate dar. Während GaAs-Substrate sehr gute Isolatoren sind, müssen Si-Substrate hoch dotiert werden, um den so genannten Latch-up-Effekt zu vermeiden, der nur bei komplementären Techniken, hierfür steht das C in CMOS, also in Technologien mit p- und n-Kanal-Transistoren auftritt. Diese Substrate verursachen nun aber bei Spulen und Wellenleitern erhebliche Verluste durch Wirbel- und Verschiebungsströme.

Taktrückgewinnung

Das Prinzip der Taktrückgewinnung besteht in der Synchronisation eines integrierten spannungsgesteuerten Oszillators mit den Flanken der einlaufenden Datenbits. Als Resonatoren werden dabei Schwingkreise aus Spulen und Kondensatoren verwendet. Daher war zunächst die Aufgabe zu lösen, eine integrierte Spule mit möglichst geringen Verlusten, d. h. hohem Gütefaktor zu konstruieren. Nach entsprechenden Simulationsrechnungen wurden die in Abbildung 2 gezeigten Teststrukturen auch messtechnisch im Frequenzbereich bis 10 GHz charakterisiert.

Untersucht wurde z. B. die Möglichkeit, die durch Verschiebungsströme verursachten Verluste durch einen Masseschild zwischen der Spule und dem Substrat zu reduzieren. Um zusätzliche Wirbelströme in diesem aus Metall oder polykristallinem Si bestehenden Schild zu vermeiden, muss er in geeigneter Weise geschlitzt ausgeführt werden. Als noch effektiver erwies sich jedoch die Nutzung der zwei obersten der insgesamt 6 Metallisierungslagen, die zudem doppelt so dick wie die 4 darunter

liegenden sind, um zwei Windungen übereinander anzuordnen. Auf diese Weise wird einerseits die Induktivität erhöht, andererseits aber auch die Fläche und damit die Wechselwirkung mit dem Substrat verringert und schließlich sogar der ebenfalls sehr entscheidende Ohmsche Widerstand der Spule verringert. Somit konnten bei 10 GHz ein Gütewert von 6,1, eine Induktivität von 0,6 nH und eine Selbstresonanzfrequenz von 40 GHz erreicht werden. Als weiteres Grundbauelement wird eine abstimmbare, genauer gesagt eine spannungsabhängige Kapazität, auch Varaktor genannt, benötigt. Die CMOS-Technologie bietet hier eine Vielzahl von Möglichkeiten: So könnte man den Übergang der Drain- und Source-Kontakte von n-Kanaltransistoren zur p-Wanne ausnutzen. Dabei ergeben sich zwar die höchsten Flächenkapazitätswerte, jedoch ist die Abstimmung sehr nichtlinear und durch die Flussspannung des pn-Übergangs von ca. 0.6 V begrenzt. Dieses Problem tritt bei der Gatekapazität nicht auf, wobei lediglich n-leitende Kanäle sinnvoll sind, da die Elektronen gegenüber den Löchern eine erheblich größere Beweglichkeit aufweisen und so serielle Widerstände gering gehalten werden können. In einem gewöhnlichen n-Transistor wird der

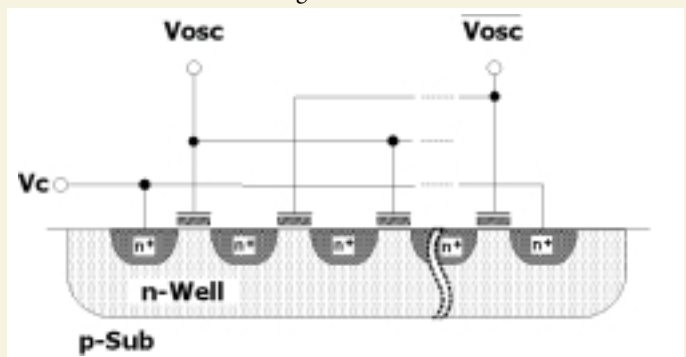


Abb. 3: Varaktor mit n-Akkumulationsschicht.

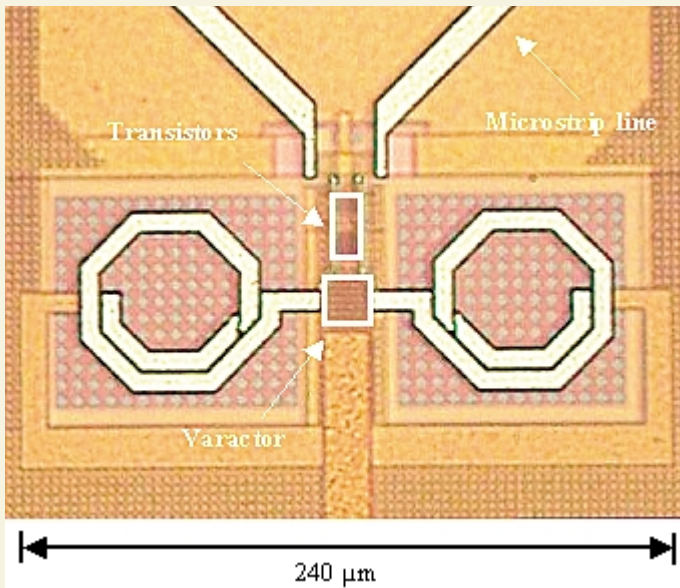


Abb. 4: LC-Oszillator

Kanal durch Inversion an der Oberfläche einer p-Wanne erzeugt. Dieser Prozess ist jedoch relativ langsam und zudem wiederum von der niedrigen Beweglichkeit der Löcher beeinflusst. Günstiger sind daher Akkumulationsgebiete in einer n-Wanne, wie in Abbildung 3 dargestellt.

Auf diese Weise wurden Gütefaktoren von 18-29, eine Flächenkapazität $0,5 \text{ fF}/\mu\text{m}^2$ und ein Abstimmverhältnis von 1,7 erreicht. Auf der Grundlage dieser Bauelemente sowie eines rückgekoppelten Differenzverstärkers konnte nun ein spannungsgesteuerter 10 GHz LC-Oszillator [2] aufgebaut werden, der in Abbildung 4 dargestellt ist. Messungen ergaben einen Abstimmungsbereich von mehr als 1 GHz und eine Verlustleistung von lediglich

7,2 mW. Die Frequenzstabilität eines Oszillators wird durch sein Phasenrauschen beschrieben. Hier wurde ein Wert von -105 dBc/Hz bei einem Abstand von 1 MHz zur Mittenfrequenz gemessen, d. h. die spektrale Leistungsdichte ist in 1 MHz Abstand vom gewünschten Signal bereits um mehr als 10 Zehnerpotenzen abgesunken. Um die Güte der Gesamtanordnung weiter zu verbessern, wurde außerdem eine Kaskadierung von 4 mit LC-Resonanzkreisen belasteten Differenzverstärkern zu einem Ringoszillator untersucht. Obwohl hierbei die Güte etwa mit der Zahl der Stufen linear zunimmt, konnte das Phasenrauschen insgesamt nicht verbessert werden, da die Stromquellen zusätzliche Rauschleistung erzeugen und den gewünschten Effekt somit überdecken. Daher soll auf diese künftig verzichtet werden. Der Ringoszillator mit 4 Stufen ist in Abbildung 5 gezeigt.

Zusätzlich zur Phasenregelung über die spannungsgesteuerten Kapazitäten wurde hierbei eine Injektion des differenzierten und anschließend gleichgerichteten Datensignals vorgesehen, so dass der Oszillator also bei jedem Wechsel der Eingangsdaten einen Impuls erhält und somit zusätzlich – ähnlich einer regelmäßig angestoßenen Kinderschaukel – synchronisiert wird.

Datenentscheider

Zu den durch den Takt nunmehr präzise vorgegebenen Zeitpunkten muss jetzt entschieden werden, ob es sich bei dem jeweiligen Bit um eine 0 oder eine 1 handelt. Obwohl grundsätzlich die Möglichkeit der Wiederherstellung von Daten der wesentlichste Vorteil der Digital- gegenüber der Analogtechnik ist – man vergleiche nur die Klangqualität einer Langspielplatte mit der einer CD – können selbst bei dieser groben Entscheidung aufgrund der großen zu überbrückenden Distanzen Fehler

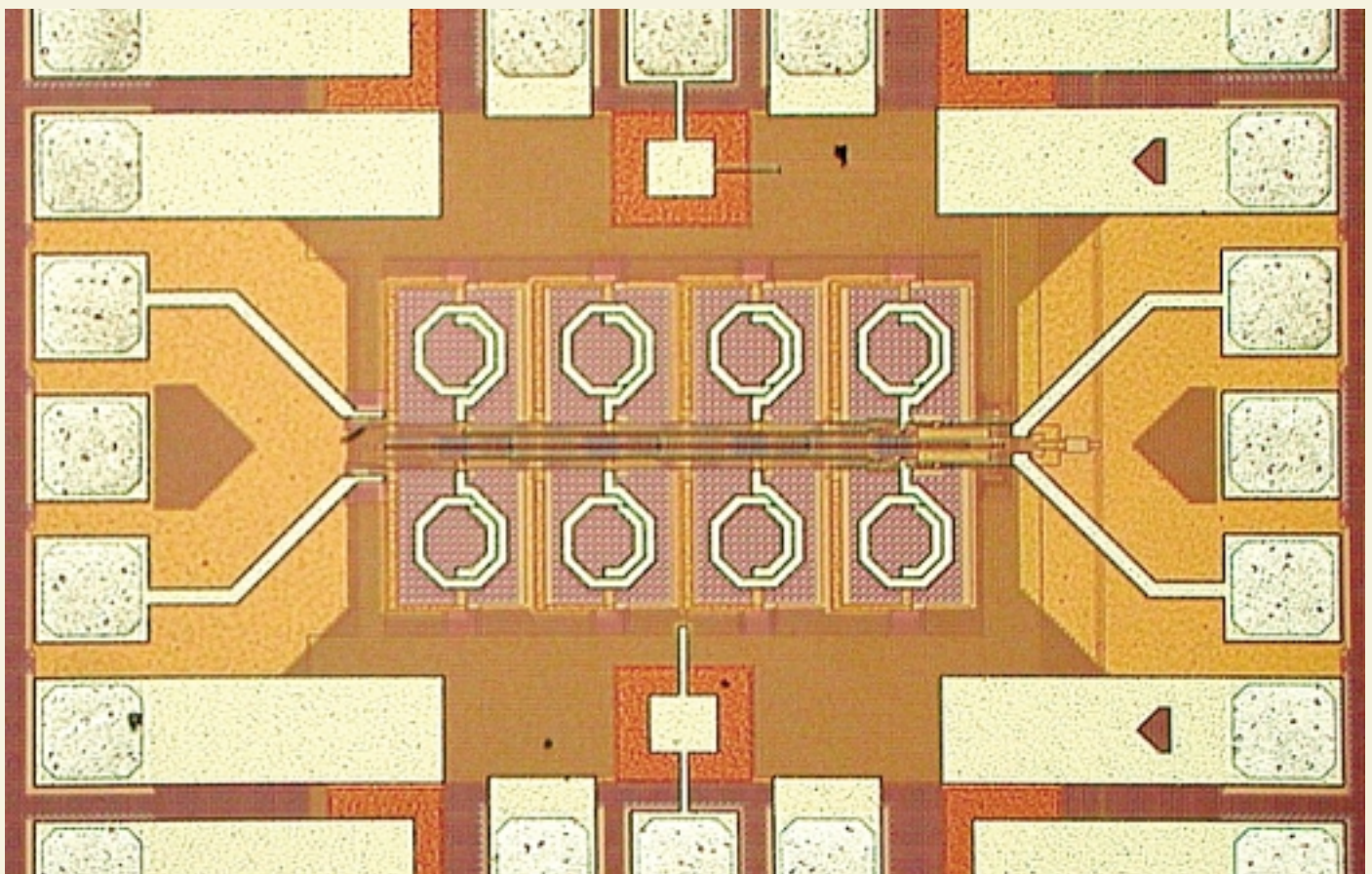


Abb. 5: 4-stufiger LC-Ringoszillator mit Injektion.

in erheblicher Zahl auftreten. Zunächst ist zu berücksichtigen, dass verschiedene Rauschprozesse auf den Kanal einwirken, die jedoch zufällig und daher nicht kompensierbar sind. Ferner treten so genannte Intersymbolinterferenzen auf. Darunter versteht man die Beeinflussung des Signals zu einem bestimmten Zeitpunkt durch das zuvor übertragene Bit. Dieser Effekt hat seine Ursache in der endlichen Bandbreite der den Übertragungskanal bildenden Komponenten und in der sowohl von seiner Wellenlänge (chromatische Dispersion) als auch der Polarisierung (Polarisationsmodendispersion) abhängigen Ausbreitungsgeschwindigkeit des Lichtes in der Glasfaser. Vor allem der letztgenannte Effekt ist sehr störend, da er zeitvariant ist und Zeitkonstanten im Millisekundenbereich – hervorgerufen etwa durch ein über die Faser rollendes Fahrzeug – aber auch im Bereich von Tagen – zum Beispiel durch jahreszeitlich bedingte Temperaturveränderungen – haben kann. Hier sind also adaptive Kompensationsverfahren erforderlich, wobei elektronische Lösungen gegenüber den auch in Paderborn entwickelten optischen Methoden deutlich kostengünstiger sind, da sie in den Empfänger nahezu ohne zusätzlichen Aufwand integrierbar sind. Grundsätzlich zu unterscheiden sind dabei lineare Transversalfilter und nichtlineare Rückkopplungsschleifen. In einem Transversalfilter werden mehrfach verzögerte Kopien des Eingangssignals mit Faktoren gewichtet und aufsummiert. Die Gewichtungsfaktoren werden ständig so angepasst, dass die Intersymbolinterferenzen optimal kompensiert werden. Alternativ kann die Kenntnis über das zuletzt erkannte Bit genutzt werden, um mit Hilfe einer geeigneten Rückkopplung die Entscheidung über das nachfolgende Bit zu beeinflussen. Abbildung 6 zeigt eine Realisierung dieses Ansatzes. Zunächst wird das links einlaufende Signal auf zwei Kanäle verteilt, die mit dem Taktsignal und dem invertierten Taktsignal jeweils halber Frequenz arbeiten, so dass die Geschwindigkeitsanforderungen reduziert werden können. Anschließend wird das aktuelle Bit in jedem Kanal mit zwei unterschiedlichen, ebenfalls adaptiven Referenzspannungen verglichen, die den beiden Fällen entsprechen, dass zuvor eine 1 bzw. eine 0 empfangen wurden. Sobald diese Entscheidung im jeweils anderen Kanal vorliegt, wird das zugehörige Entscheidungsergebnis mit Hilfe eines Selektors ausgewählt, an den Ausgang weitergeleitet und gleichzeitig nun seinerseits verwendet, um die richtige Entscheidung über das nachfolgende Bit im wiederum anderen Kanal auszuwählen.

Dieses Schaltungsprinzip wurde gemeinsam mit einer Gruppe der Bell Laboratories in den USA entwickelt und zum Patent angemeldet. Der Test erfolgte zunächst in einer GaAs-Technologie, wobei eine maximale Datenrate von 20 Gbit/s erreicht wurde [3]. Abbildung 7 zeigt ein Chipfoto: Auf dem insgesamt 2 mm x 3 mm großen Schaltkreis sind aufgrund der in dieser Technologie wesentlich größeren Verdrahtung deutlich die regelmäßigen Strukturen von Daten-Flip-Flops erkennbar. Die Zuführung der schnellen Takt- und Datensignale erfolgt mit differentiellen koplanaren Wellenleitern. Nunmehr wird eine entsprechende Lösung in Si-CMOS-Technologie erarbeitet, die zusätzlich mit einem Transversalfilter kombiniert werden soll. Hierfür wurden bereits spezielle Daten-Flip-Flops entwickelt, die auch bei den maximal zulässigen Versorgungsspannungen von nur 1.8 V zuverlässig arbeiten. Unter Beibehaltung der sehr schnellen und weniger störanfälligen differentiellen Stromschalterlogik wurde auf die Konstantstromquelle verzichtet. Zur Charakterisierung der Flip-Flops wurden sie zu statischen

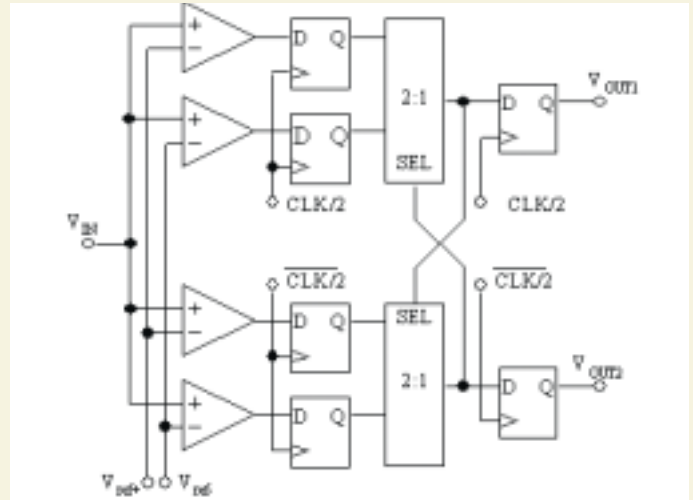


Abb. 6: Schaltbild des Datenentscheiders.

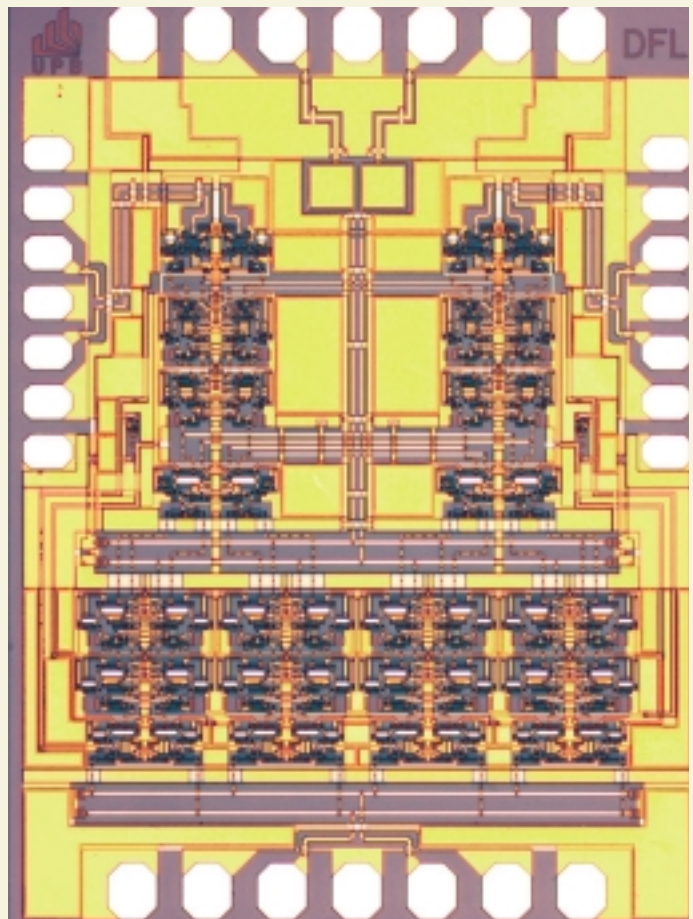


Abb. 7: Chipfoto des Datenentscheiders.

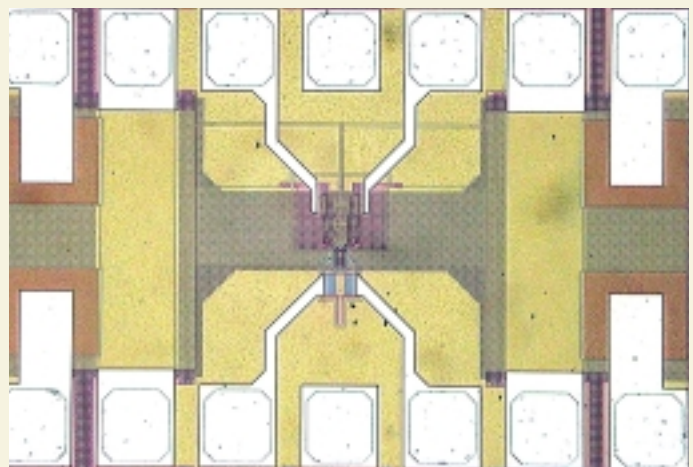


Abb. 8: Chipfoto des Frequenzteilers.

Frequenzteilern verschaltet, mit denen Eingangsfrequenzen bis zu 18 GHz [4] verarbeitet werden konnten. Abbildung 8 zeigt wiederum ein Chipfoto. Die Größe dieses nur 0.7 mm x 0.5 mm messenden Chips ist allein durch die oben und unten angebrachten Kontaktreihen bestimmt. Die gesamte Schaltung konzentriert sich in der Chipmitte. Die einzelnen Flip-Flops sind hier nicht erkennbar.

Geschwindigkeitsgrenzen noch immer nicht in Sicht

Doch auch mit 10 Gbit/s sind die Geschwindigkeitsgrenzen der Si-CMOS-Technologie nicht ausgereizt. Schon heute haben führende Hersteller mit 0.12 μm - und 0.09 μm -Technologien erste digitale Schaltungen bis zu Datenraten von 40 Gbit/s und Oszillatoren bis zu Frequenzen von knapp 100 GHz vorgestellt, und auch 0.03 μm -Technologien sind bereits absehbar. Auf germaniumhaltige Schichten aufgebraute Kanäle erhöhen die Geschwindigkeit der Elektronen in den Transistoren zusätzlich. Die hiermit verbundenen neuen schaltungstechnischen Möglichkeiten aber auch Probleme werden der Gegenstand nicht weniger interessanter Forschungsprojekte in den kommenden Jahren sein.

Literatur

- [1] A. Thiede, Z.-G. Wang, M. Schlechtweg, M. Lang, P. Leber, Z. Lao, U. Nowotny, V. Hurm, M. Rieger-Motzer, M. Ludwig, M. Sedler, K. Köhler, W. Bronner, J. Hornung, A. Hülsmann, G.

Kaufel, B. Raynor, J. Schneider, T. Jakobus, J. Schroth, M. Berroth: „Mixed Signal Integrated Circuits Based on GaAs HEMT's“, IEEE-Trans. VLSI Systems, vol.6(1998), no.1, pp.6-17.

- [2] Z. Gu, B. Bartsch, A. Thiede, R. Tao, Z.-G. Wang: „Fully Integrated 10 GHz CMOS LC VCOs“, European Microwave Conference, Munich/Germany, 2003, pp.583-586.
- [3] L. Möller, Z. Gu, A. Thiede, S. Chandrasekhar, L. Stulz: „20 Gb/s electrical data recovery using a decision feedback equalizer supported receiver“, Electronics Letters, vol. 39(2003), no.1, pp.78-79.
- [4] Z. Gu, A. Thiede: 18 GHz Low-Power CMOS Static Frequency Divider Electronics Letters, vol. 39(2003), no.20, pp.1433-1434.



M.Sc. Zheng Gu studierte an der Südostuniversität Nanjing, China. Seit Januar 2001 arbeitet er als Doktorand am Fachgebiet Höchstfrequenzelektronik. Im Rahmen eines seit Februar 2002 von der DFG geförderten Projektes entwickelte er integrierte Schaltkreise auf der Basis einer industriellen 0.18 μm CMOS-Technologie für die Datenrückgewinnung bei einer Datenrate von 10 Gbit/s.

...und nach
den Vorlesungen
Entspannung pur
in der
DRIBURG THERME!



...wo der Mensch sich wohl fühlt...

33014 Bad Driburg • Tel. 0 52 53 / 70 116

www.driburg-therme.de

info@driburg-therme.de

GUTSCHEIN

~~10,50 €~~
> 8,50 €

Bei Vorlage dieses Gutscheins erhalten Sie eine zusätzliche Ermäßigung auf den Tagestarif Bad + Sauna für Studenten von **2,00 €**.

~~10,50 €~~
> 8,50 €

Bei Vorlage dieses Gutscheins erhalten Sie eine zusätzliche Ermäßigung auf den Tagestarif Bad + Sauna für Studenten von **2,00 €**.

Ganz schön scharf!

ORGA Systems bringt die richtige Würze in die mobile Telekommunikation - weltweit

**ORGA
Systems**

Neugierig geworden?

Unser Erfolgsrezept findet man unter www.orga-systems.com



ORGA Systems | Am Hoppenhof 33 | D-33104 Paderborn
Tel.: + 49 (0) 52 51-889 0 | Fax: + 49 (0) 52 51-889 37 37

www.orga-systems.com

die
Sprach-
werkstatt

Unser Können - Ihre Chance

- EDV
- Fremdsprachen
- Deutsch als Fremdsprache
- Kaufmännische Seminare
- Umschulungen
- Prüfungsvorbereitung



**Informieren Sie sich in
unserem Beratungszentrum!**

Mo - Do: 07.30 - 20.00 Uhr
Fr: 07.30 - 16.00 Uhr



Privates Institut
für Kommunikation,
Wirtschaft und
Sprache GmbH
Stettiner Straße 40-42
33106 Paderborn

Tel. 05251/77999-0
Fax 05251/77999-79
www.die-sprachwerkstatt.de
paderborn@die-sprachwerkstatt.de



jetzt neu:

Tickets für alle großen Events

- Rock- und Pop-Konzerte
- Musical, Oper, Theater
- Sport - Fussball - Formel 1
- und alles andere

Lassen Sie sich umfassend bei uns beraten!

**Universitätsbuchhandlung Meier KG
Warburger Str. 98 - 33098 Paderborn**

Ticket-Hotline: 05251 - 180590

Rationalismus zwischen Vision und Wirklichkeit

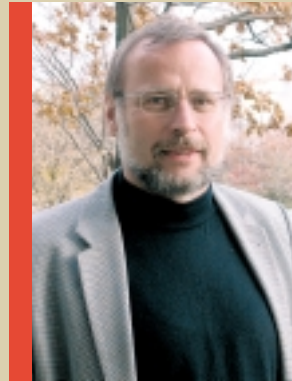
Descartes, Leibniz und die Grenzen des Machbaren

Hat der Rationalismus, der doch so entscheidend Wissenschaft und Technik in unserer Welt geprägt hat, abgewirtschaftet? Zwingt uns nicht die Einsicht in die Grenzen der Vernunft, unsere alltägliche Erfahrung von Differenz und Unsicherheit in die Schranken? Nicht unbedingt, denn der Rationalismus bietet Strategien, mit dieser Erfahrung umzugehen. Der Lehrstuhl für Wissenschaftstheorie und Philosophie der Technik hat sich zur Aufgabe gesetzt, diese Strategien, die dabei eingesetzten Methoden, deren Reichweite und Grenzen zu erforschen. Im weiten Feld des Nachdenkens über Wissenschaft und Technik geht es ihm vor allem um die logischen und kognitiven Bedingungen des Erkennens, des wissenschaftlichen und technischen Handelns.

Mit wissenschaftshistorischer Zugangsweise und einem Forschungsschwerpunkt in der Geschichte der neueren Logik kann dabei die logikhistorische Tradition Paderborns, die mit den Namen Franz Schupp, Christoph Kann und Stephanie Weber-Schroth verbunden ist, fortgesetzt werden. Im Rahmen des vom Deutschen Akademischen Austauschdienst geförderten Projektes „Formalsprachen als Universalsprachen und die Ursprünge der modernen Logik“, an dem Wissenschaftler aus Buenos Aires und La Plata in Argentinien beteiligt sind, wird u. a. eine Neubewertung der Methoden des Rationalismus versucht.

Vom Scheitern des Rationalismus

Der Rationalismus der Neuzeit, wie er sich seit René Descartes (1596-1650) und Gottfried Wilhelm Leibniz (1646-1716) entwickelte, wird üblicherweise mit dem Bestreben gleichgesetzt, das mögliche Wissen unter eine allumfassende Einheit zu bringen. Seine Orientierung am Exaktheitsideal der Mathematik wird hervorgehoben, wie auch sein Programm, eine *Mathesis universalis* zu begründen, die ein Rechnen auch außerhalb der Mathematik der Mathematiker ermöglicht. Der Prozess der Erkenntnisgewinnung wird zu einem quasi-automatischen Verfahren. Das Wissen wird absolut, jede Unsicherheit soll aufgehoben werden. In diesem Sinne hat schon Edmund Husserl in seiner Krisis-Schrift den Cartesischen Rationalismus gesehen, wenn er mit Descartes die völlig neue Idee entstehen lässt, „dass die unendliche Allheit des überhaupt Seienden in sich eine rationale Alleinheit sei, die korrelativ durch eine universale Wissenschaft, und zwar restlos zu beherrschen sei“ (Husserl 1976, 20). Es ist genau diese Idee, die Max Horkheimer und Theodor W. Adorno als Kennzeichen der Aufklärung angesehen haben, einer Aufklärung allerdings, die letztlich in die Katastrophe des tota-



Prof. Dr. phil. Volker Peckhaus wurde 2002 auf den Lehrstuhl für Wissenschaftstheorie und Philosophie der Technik an der Fakultät für Kulturwissenschaften der Universität Paderborn berufen. Er ist Mitglied des Vorstandes des Heinz-Nixdorf-Instituts. Zu seinen Arbeitsgebieten gehören u. a. Geschichte und Philosophie der Logik und der mathematischen Grundlagenforschung. Die am Lehrstuhl betreute bio-bibliographische *Database for the History of Logic* steht Forschern offen.

len Staates geführt habe. Aufklärung ist für sie „totalitär wie nur irgendein System“ (Horkheimer/Adorno 1969, 31). Ihre Vorbedingung ist die formale Logik, „die große Schule der Vereinheitlichung“, die „das Schema der Berechenbarkeit der Welt“ geboten habe (ebd., 13), und die Methode des Totalen ist das mathematische Verfahren (ebd., 31):

Wenn im mathematischen Verfahren das Unbekannte zum Unbekannten einer Gleichung wird, ist es damit zum Altbekannten gestempelt, ehe noch ein Wert eingesetzt ist. Natur ist, vor und nach der Quantentheorie, das mathematisch zu Erfassende; selbst was nicht eingeht, Unauflöslichkeit und Irrationalität, wird von mathematischen Theoremen umstellt. In der vorwegnehmenden Identifikation der zu Ende gedachten mathematisierten Welt mit der Wahrheit meint Aufklärung, vor der Rückkehr des Mythischen sicher zu sein. Sie setzt Denken und Mathematik in eins. Dadurch wird diese gleichsam losgelassen, zur absoluten Instanz gemacht.

Damit ist aber ein neuer Fetisch gegeben. In einer dialektischen Bewegung verfällt die Aufklärung wieder dem Mythischen, vor dem sie sich sicher wähnte. Die postmoderne Kritik geht noch weiter, indem sie mit Jean-François Lyotard die Aufklärung als Metaerzählung zur Legitimation des Wissens für gescheitert erklärt, gescheitert durch Zerstreuung der Sprachspiele und Zersplitterung in unterschiedlichste Wissenschaftssprachen, Logiken und Codes (Lyotard 1999, §§ 9-10).

Ist aber eine solche Kennzeichnung des Rationalismus und die daraus abgeleitete Kritik gerechtfertigt? In diesem Beitrag soll an einem historischen Beispiel gezeigt werden, dass die frühen Rationalisten durchaus zwischen Vision und Machbarkeit zu unterscheiden wussten und Strategien besaßen, Utopien, wie sie im Gedanken der Allumfassendheit repräsentiert sind, handhab-

bar zu machen. An diesen Strategien und Verfahren, dem visionär Angezielten näher zu kommen, sollte der Rationalismus gemessen werden.

Wissenschaftliche Universalsprache

Im Jahre 1629 schickte der Pater Marin Mersenne an René Descartes (Abbildung 1) einen Brief, in dem er von dem Projekt einer „nouvelle langue“ eines gewissen Vallée berichtete, der eine „langue matrice“ gefunden haben wollte, die es ihm angeblich erlaubte, alle Sprachen zu verstehen. In seinem berühmten Antwortschreiben vom 20. November 1629 wiederholte Descartes zunächst bekannte Argumente über Möglichkeiten und Probleme von Pasigraphien (Allgemeinschriften), Polygraphien (Schriften für mehrere Sprachen) und Steganographien (Geheimschriften). Die Grammatik solcher Universalsprachen, so meinte er, muss einfach und regelmäßig sein, darüber hinaus muss ein vollständiges System der elementaren Begriffe aufgestellt werden. Jeder Begriff kann dann mit einer Nummer versehen werden, die als Schlüssel für die Zuweisung von Synonymen anderer Sprachen dient. Ein solches Vorgehen funktioniert natürlich nur in der schriftlichen Kommunikation. Wenn einer die fremde Sprache auch sprechen will, muss er zusätzlich den ganzen Wortschatz dieser Sprache erlernen.

Descartes führt diese Überlegungen aber noch weiter. Um die Elementarbegriffe nicht nur erlernen, sondern auch behalten zu können, müssen sie wie Ideen und Gedanken geordnet werden. Diese Ordnung sollte aber der Ordnung der Zahlen entsprechen, denn letztere brauchen ja nicht einzeln erlernt, sondern können durch Reihung erzeugt werden. Die Schaffung einer universellen Sprache hängt also von der Schaffung einer wahren Philosophie ab, in der alle *einfachen* Ideen benannt und bezeichnet werden und in der dann durch Rechnung alle nur *denkbaren* klaren und deutlichen Ideen erzeugt werden können. Für Descartes ist das der bestmögliche Kunstgriff, um zu einer guten Wissenschaft zu kommen. Descartes bleibt jedoch skeptisch hinsichtlich der Durchführbarkeit eines solchen Programms, denn er schließt den Brief mit folgendem Ausblick (zitiert nach Eco 1997, 225f, Descartes 1987, 81f):

Nun glaube ich zwar, dass solch eine Sprache möglich ist und dass man die Wissenschaft finden kann, von der sie abhängt und mittels derer die Bauern dann besser werden über die Wahrheit urteilen können, als es heutzutage die Philosophen tun. Aber ich kann mir nicht vorstellen, wie sie jemals in Gebrauch kommen soll: Sie setzt große Veränderungen in der Ordnung der Dinge voraus, und es müsste erst die ganze Welt ein irdisches Paradies werden, was man nur im Land der Romane erwarten kann.

Descartes formuliert hier den Gedanken einer philosophischen oder rationalen Sprache, die als Ideographie das System der menschlichen Gedanken vollständig abbildet, indem sie die Behauptung der (logischen) Möglichkeit, eine vollständige Liste der elementaren Ideen und der damit korrespondierenden elementaren Begriffe angeben zu können, mit einer *mathesis universalis* verbindet, mit deren Hilfe dann alles nur Denkbare rechnerisch konstruiert werden kann. Die oben zitierte Stelle zeigt aber auch, dass Descartes selbst offenbar nicht gewillt war, die Probleme bei der Formulierung einer solchen Sprache „frontal anzupacken“, wie es Umberto Eco in seinem Buch *Die Suche*



Abb. 1: René Descartes (1596-1650), Gemälde von Frans Hals, Statens Museum for Kunst, Kopenhagen.

nach der vollkommenen Sprache formuliert hat (Eco 1997, 226). Obwohl Descartes von der logischen Möglichkeit einer universalen Sprache und damit wahren Philosophie überzeugt war, ging er doch zugleich von der praktischen Utopie dieses Gedankens aus. Die allumfassende philosophische Sprache wird zur Vision, auf die die wissenschaftliche Praxis abzielt, die zugleich aber für den Menschen als endliches Wesen in unendliche Ferne rückt. Dies mag Descartes davon abgehalten haben, ein solches Programm zielstrebig anzugehen.

Ein ganz anderer Zugang findet sich bei Leibniz (Abbildung 2), der die wissenschaftliche Universalsprache als Hilfsmittel für den kontrollierten Wissenserwerb eingesetzt sehen wollte. Leibniz operationalisierte den von Descartes als utopisch eingeschätzten Gedanken einer philosophischen oder rationalen Sprache. Im Leibnizschen Nachlass findet sich eine auszugsweise Abschrift des Briefes von Descartes an Mersenne, versehen mit einem Kommentar von Leibniz' Hand (C, 27-28). Selbst wenn die von Descartes vorgeschlagene Sprache von einer wahren Philosophie abhinge, so schreibt Leibniz dort, so hinge sie jedoch damit noch nicht von deren Perfektion ab. Man könne eine solche Sprache einrichten, auch wenn die Philosophie noch nicht perfekt sei. In dem Maße, in dem sich die Wissenschaft des Menschen weiterentwickle, würde sich auch die Sprache weiterentwickeln.

Mit der Organisation des Wissens sollte nach Leibniz' Auffassung also begonnen werden, auch wenn das Organisationsmittel, die philosophische Sprache, noch nicht vollständig vorliegt. Die Vollendung einer solchen philosophischen Sprache, in der alle einfachen Ideen formuliert, alle möglichen Ideen aber formulierbar sind, wäre im Rahmen Leibnizscher Metaphysik ohnehin nicht zu erwarten. Im unendlich komplexen System der prästabilierten Harmonie ist es dem Menschen nicht möglich, den vollen Zugriff auf die im Schöpfungsakt kreierten Wahrheiten zu



Abb. 2: Gottfried Wilhelm Leibniz (1646-1716), Niedersächsische Landesbibliothek, Hannover.

erlangen. Gleichwohl gilt es, methodische Hilfsmittel zu schaffen, mit denen die Reichweite des Zugriffs auf diese Wahrheiten sukzessive erweitert werden kann. Im Leibnizprogramm waren diese Aufgaben im Rahmen einer *ars inveniendi* vor allem deduktiven Verfahren wie Kombinatorik, Syllogistik und logischem Kalkül zugeordnet, die eine Wissenschaftssprache als Medium voraussetzten.

Theoria cum praxi

„Theoria cum praxi“ war das Motto, das Leibniz über sein Werk stellte und das er in seinem Brief an den Mathematiker und Naturwissenschaftler Gabriel Wagner wie folgt erläuterte (GP VII, 514-527, Zit. 525):

Die Kunst der Practick steckt darinn daß man die zufälle selbst unter das joch der wißenschafft so viel thunlich bringe. Je mehr man dieß thut, ie bequemer ist die theorie zur Practick.

Leibniz schlägt hier die theoretische Durchdringung der Praxis vor, zumindest so weit dies tunlich ist. Theorie und Praxis sind aufeinander angewiesen, ohne dass beide zusammenfallen oder die Praxis erst nach formulierter Theorie einsetzen könnte. Es liegt nahe, unter „Praxis“ das ingenieurmäßige Erfinden und Konstruieren, aber auch politisches und ökonomisches Handeln zu verstehen, Künste also, in denen sich Leibniz mit wechselndem Erfolg versucht hat. Aber seine Ausführungen gelten auch für wissenschaftliches Handeln im Allgemeinen, z. B. für das Gebiet der Sprachkonstruktion und damit eng zusammenhän-



Abb. 3: Leibniz' Medaillenenwurf zur Darstellung des binären Zahlensystems. Das Bild der Schöpfung: „Omnibus ex nihilo ducendis sufficit unum“ [„Um alles aus dem Nichts herzuleiten, genügt Eines“].

gend für Logik und Mathematik. Leibniz ließ es also zu, die Sprachkonstruktion, d. h. die Formulierung von Syntax (Grammatik) und Semantik einer Sprache, anzugehen, auch ohne dass eine vollständige Klassifikation der einfachen Ideen bereits erreicht oder auch nur erreichbar wäre. Mathematische Kalküle werden damit stets auf begrenzte Gegenstandsbereiche angewendet. Die Bestimmung dieser Gegenstandsbereiche setzt kreative Akte voraus, die sich einer vollständigen Rationalisierung sperren.

Leibniz strebte die theoretisch optimale Lösung zwar stets an, zugunsten rascher Ergebnisse ließ er aber auch Zwischenlösungen zu, die sich im Einsatz bewähren müssen und im Fortgang der Entwicklung optimiert werden können. Auch durch diesen Gedanken wird die rationale Durchdringung der Wissenschaft für nicht-rationale Kreativität geöffnet, die zum Motor der Erfindungskunst wird.

Rationalismus und Kreativität

Damit ist aber auch die landläufige Entgegensetzung von axiomatisch-deduktiver mathematischer Methode und oftmals methodenunabhängiger Erfindungstätigkeit im kreativen Prozess hinfällig. Beides wirkt zusammen.

Mit dem Argument der wechselseitigen Ergänzung kann auch die Behauptung der kreativitätstötenden Funktion des *Collegium Logicum* pariert werden, wie sie etwa den mephistophelischen Ratschlägen an den Scholaren in Goethes *Faust* zu entnehmen ist (*Faust I*, 1911-1917). Der Scholar wünscht sich zwar, recht gelehrt zu werden und alles, was auf Erden und im Himmel ist, zu erfassen, die Wissenschaft und die Natur zu studieren, gleichwohl möchte er an schönen Sommertagen auch Freiheit und Zeitvertreib genießen. Davon rät Mephistopheles ab:

Gebraucht der Zeit, sie geht so schnell von hinnen,/ Doch Ordnung lehrt Euch Zeit gewinnen./ Mein teurer Freund, ich rat Euch drum/ Zuerst Collegium Logicum./ Da wird der Geist Euch wohl dressiert,/ In spanische Stiefeln eingeschnürt,/ Daß er bedächtiger so fortan/ Hinschleiche die Gedankenbahn,/ Und nicht etwa, die Kreuz und Quer',/ Irrlichtere hin und her.

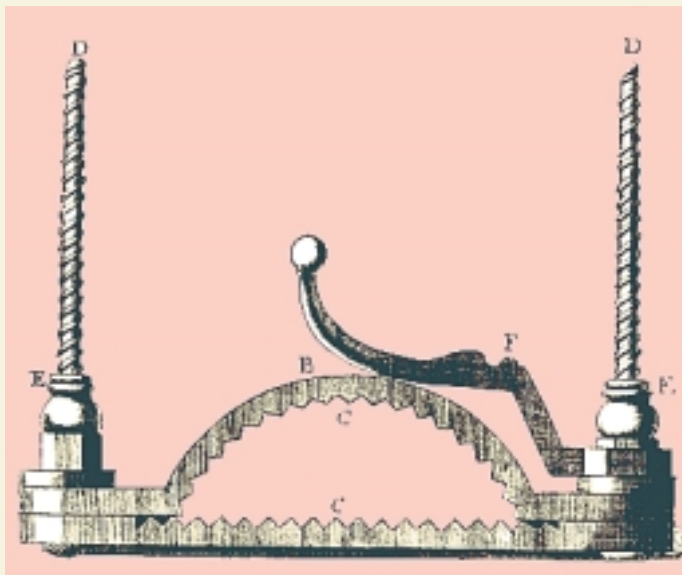


Abb. 4: Spanische Stiefel, aus der Constitutio Criminalis Theresiana, 1769.

Die Logik wird hier zum Folterinstrument des Geistes stilisiert. Eine ganz ähnliche Haltung gegenüber mathematischen oder logischen Automatismen scheint auch Georg Christoph Lichtenberg (1742-1799) gehabt zu haben, einer der genialen und kreativen Geister des 18. Jahrhunderts, der gleich im ersten Heft seiner *Sudelbücher* schrieb (1998, A 11):

Die Erfindung der wichtigsten Wahrheiten hängt von einer feinen Abstraktion ab, und unser gemeines Leben ist eine beständige Bestrebung uns zu derselben unfähig zu machen, alle Fertigkeiten, Angewohnheiten, Routine, bei einem mehr, als bei dem andern, und die Beschäftigung der Philosophen ist es, diese kleinen blinden Fertigkeiten, die wir durch Beobachtungen von Kindheit an uns erworben haben, wieder zu verlernen. Ein Philosoph sollte also billig als ein Kind schon besonders erzogen werden.

Umso überraschender ist es, dass Lichtenberg zu diesen Zeilen von Leibniz inspiriert worden war. Er begann mit seinen Aufzeichnungen, kurz nachdem die Leibnizschen *Nouveaux Essais* aus dem Nachlass veröffentlicht worden waren und große Aufmerksamkeit in der gelehrten Welt erregt hatten. Lichtenberg nahm Leibniz als Beispiel für diesen besonderen Typus Mensch, den er „Philosoph“ nannte. Er zitierte das Leibnizsche Bekenntnis, in allen Wissenschaften, die er gelernt habe, gleich erfinden wollen, auch wenn er deren Prinzipien noch gar nicht besessen habe. Dies habe Leibniz dazu bewogen, auf die Grundlagen der Wissenschaften zurückzugehen und sich in allen Fällen durch eigene Regeln aus den Problemen herauszuhelfen (ebd., A 12). Lichtenberg erwähnte auch Leibniz' Vision von der Möglichkeit, das *alphabetum cogitationum humanarum* zu finden, und dann durch Kombination der darin enthaltenen Buchstaben und durch Analyse der damit gebildeten Worte alles zu erfinden und zu beurteilen. Der Gedanke der Dienstbarmachung der Logik als Organon im Rahmen einer *ars inveniendi* erlaubte es für Lichtenberg, vermeintliche Routine und Fruchtlosigkeit einer lediglich als „Prüfstein“, wie Kant dies später wollte, eingesetzten Logik zu durchbrechen. Rationalität und Kreativität müssen sich daher nicht ausschließen, so zumindest die Meinung eines großen Kreativen.

Rationalismus und Heuristik

Die Leibnizsche Haltung kann als paradigmatisch für dynamische Theoriebildung angesehen werden, die sich auf das Machbare konzentriert, ohne allerdings die (wohl utopische) Zielvorstellung eines allumfassenden Systems vollständig aufzugeben. Einheitlichkeit und Universalität haben somit allenfalls eine heuristische Funktion. Sie wirken als regulative Ideen, geben dem Streben eine Richtung. Aber ist es nicht irrational, sich das Unerreichbare zum Ziel zu setzen? Das Begehen eines Weges mit unerreichbarem Ziel ist dann sinnvoll, wenn sich auf diesem Weg ein Nutzen für den Menschen ergibt. In der aufklärerisch-humanistischen Attitude eines Descartes-Leibnizschen Rationalismus ist das Absolute nicht alleinige Richtschnur für das Handeln des Wissenschaftlers. Der Rationalismus strebt auch die Verbesserung der Stellung des Menschen gegenüber der ihn umgebenden und sich ihm widersetzenden Natur an. Der Rationalismus ist progressiv und optimistisch, weil er davon ausgeht, dass es möglich ist, den menschlichen Zugriff auf das ihm bis dato Verschllossene zu erweitern, damit aber auch seinen Anteil am Göttlichen zu vergrößern.

Unzweifelhaft hat die kritische Sicht auf Philosophie und Wissenschaften seit Kant unsere Einsicht in die Grenzen der Vernunft und der Kalküle erweitert. Im auch theologisch motivierten Weltbild des Rationalismus hätte es aber nicht des mathematischen Aufweises unentscheidbarer Sätze, nicht überschreitbarer Komplexitätsschranken oder quantenphysikalischer Unschärferelationen bedurft, um die Existenz solcher Grenzen nachzuweisen. Für Descartes und Leibniz ergeben sich diese Grenzen schon aus der nichthintergehbaren Differenz zwischen dem endlichen Menschen und einem in jeder Hinsicht vollkommenen Schöpfer der Welt, in der wir leben.

Literatur

- Descartes, René: *CŒuvres de Descartes. Correspondance I*. Avril 1622-Février 1638, hg. Charles Adam und Adam Tannery, nouvelle présentation J. Vrin: Paris 1987.
- Eco, Umberto: *Die Suche nach der vollkommenen Sprache*, aus dem Italienischen von Burkhart Kroeber, Deutscher Taschenbuch Verlag: München 1997; Originalausgabe: *La ricerca della lingua perfetta nella cultura europea*, Laterza: Roma/Bari 1993.
- Horkheimer, Max und Theodor W. Adorno: *Dialektik der Aufklärung. Philosophische Fragmente*, S. Fischer: Frankfurt a. M. 1969.
- Husserl, Edmund: *Die Krisis der europäischen Wissenschaften und die transzendente Phänomenologie*, hg. v. W. Biemel, 2. Aufl., The Hague 1976 (Husserliana VI).
- Leibniz, Gottfried Wilhelm: *Die philosophischen Schriften von Gottfried Wilhelm Leibniz*, hg. v. C. I. Gerhardt, 7 Bde., Weidmannsche Buchhandlung: Berlin 1875–1890 [zitiert mit „GP“].
- Leibniz, Gottfried Wilhelm: *Opusculs et fragments inédits de Leibniz. Extraits des manuscrits de la Bibliothèque royale de Hanovre*, hg. v. L. Couturat, Alcan: Paris 1903 [zitiert mit „C“].
- Lichtenberg, Georg Christoph: *Schriften und Briefe*, Bd. 1: *Sudelbücher I*, 6. Aufl., Zweitausendeins: Frankfurt a. M. 1998.
- Lyotard, Jean-François: *Das postmoderne Wissen. Ein Bericht*, hg. v. P. Engelmann, Edition Passagen: Wien 1999.

Mobile Ad-hoc-Netzwerke

Das W-LAN der Zukunft

Die Geschichte der Informatik und der Elektrotechnik wird bestimmt durch die rasante Weiterentwicklung von Hardware. Immer schnellere, leistungsfähigere und kleinere Prozessoren bestimmen den Fortschritt in der Informatik. Eine ähnliche Entwicklung betrifft Sender und Empfänger in Funkgeräten und ermöglicht dort die Verbesserung der Übertragungssicherheit und der Übertragungsgeschwindigkeit.

Manchmal wird der technologische Fortschritt nicht durch neue Hardware vorangetrieben, wie ein schnellerer Prozessor oder neue Antennen, sondern durch bessere Software, also Algorithmen. Ein Beispiel hierzu sind mobile Ad-hoc-Netzwerke, die im Sonderforschungsbereich 376 Massive Parallelität, Teilprojekt C6 Mobile Ad-hoc-Netzwerke am Institut für Informatik an der Universität Paderborn vom Autor erforscht wurden (Abbildung 1).

Was sind mobile Ad-hoc-Netzwerke?

Ein Mobiles Ad-hoc-Netzwerk (MANET) ist ein drahtloses Kommunikationsnetzwerk, in dem alle Teilnehmer gleich berechtigt die Netzwerkorganisation wahrnehmen. In diesen Netzwerken gibt es keine bevorrechtigten Funkstationen, die besondere Aufgaben wahrnehmen, wie zum Beispiel ein Funkumsetzer auf einem Sendemast, der alleine zwischen allen Teilnehmern vermittelt (Abbildung 2). Eine Eigenschaft des mobilen Ad-hoc-Netzwerks ist die Fähigkeit der Selbstorganisation und Selbstkonfiguration.

Das Gegenstück zu mobilen Ad-hoc-Netzwerken sind die traditionell eingesetzten zentral gesteuerten Netzwerke. Für große Verbreitungsgebiete werden mehrere Zentralstationen eingesetzt, die das Einsatzgebiet in Zellen unterteilen. Deswegen wird auch die Bezeichnung „zelluläre Netzwerke“ verwendet. Solche Netzwerke sind einfacher zu steuern und zu planen als ein MANET. Auch kann man Abrechnung, Zugangskontrolle und Sicherheitsmechanismen besser kontrollieren. Sie werden weltweit eingesetzt

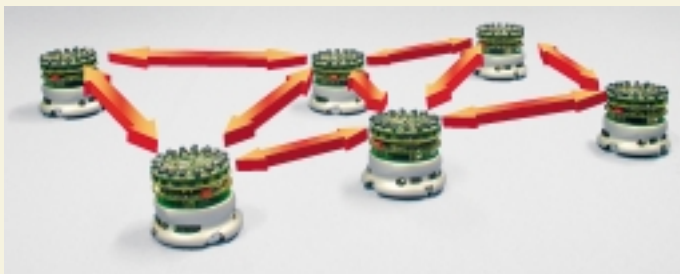


Abb. 1: Bildmontage des im SFB 376 Massive Parallelität, Teilprojekt C6, Mobile Ad-hoc-Netzwerke entwickelten Mobilen Ad-hoc-Netzwerks mit Mini-Roboter Khepera und angelegter Infrarot-Kommunikation.



PD Dr. rer. nat. Christian Schindelbauer wurde 2002 an die Universität Paderborn als Hochschuldozent für Informatik berufen. Er arbeitet in dem Fachgebiet Algorithmen und Komplexität des Instituts für Informatik an algorithmischen Problemstellungen im Bereich der Kommunikationsnetzwerke. Insbesondere arbeitet er an Algorithmen für Mobile Ad-hoc-Netzwerke, Peer-to-Peer-Netzwerke, Sensor-Netzwerke und Algorithmen für das Internet.

und ihre Abkürzungen GSM, UMTS, Bluetooth etc. sind jedem ein Begriff. Wozu benötigt man dann mobile Ad-hoc-Netzwerke?

Wozu mobile Ad-hoc-Netzwerke?

Der Hauptvorteil dieser MANETs liegt in der Einsparung des Energieverbrauchs. Jeder Mobiltelefonbenutzer kennt aus eigener Erfahrung die Unterschiede im Energieverbrauch je nachdem, ob viele Gespräche geführt worden sind oder ob nur wenige. Es macht auch einen großen Unterschied, wie weit man von der Sende-/Empfangsanlage des Mobilfunkbetreibers entfernt ist. Die notwendige Sendeleistung steigt quadratisch im Abstand, d. h. mit der Verdoppelung des Abstands wird eine Vervierfachung der Sendeenergie notwendig. Drastisch formuliert, wer doppelt so weit vom Sendemast wohnt, kann sein Gerät viermal so oft aufladen.

Dabei erscheint es widersinnig, warum zur Übermittlung einer SMS zum Nebenmann diese Nachricht zu einem weit entfernten

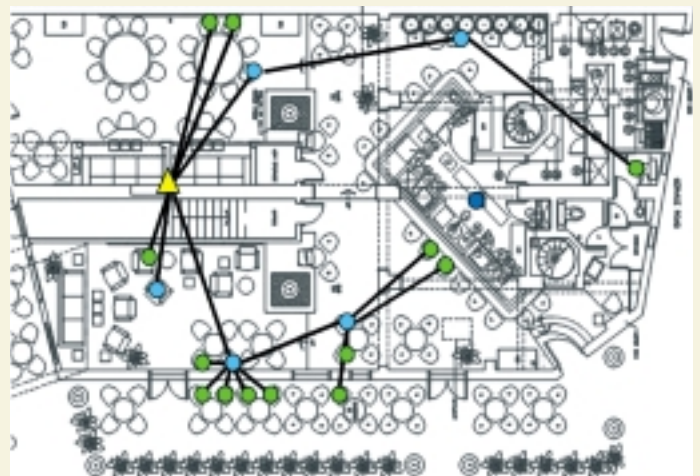


Abb. 2: Ein mobiles Ad-hoc-Netzwerk vernetzt Teilnehmer direkt untereinander.

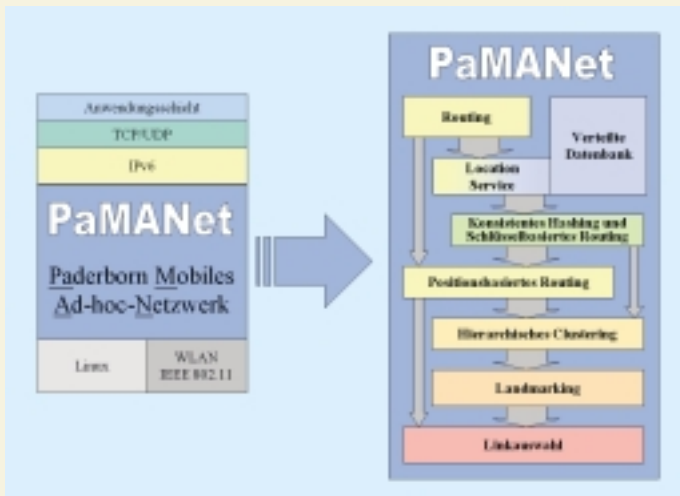


Abb. 3: Die Systemumgebung und Teilkomponenten des PaMANet (Paderborn: Mobiles Ad-hoc-Netzwerk).

Sendemast gelangen muss und von dort wieder zurückgesendet wird. Schließlich verfügen beide Geräte über Sende- und Empfangseinheiten und es gibt keinen logischen Grund, warum diese Geräte nicht direkt kommunizieren können.

Stellen wir uns ein Netzwerk für Mobiltelefone vor ohne Mobilfunkbetreiber und Sendemasten. Jedes Gerät würde nur bis zum nächsten Nachbar senden müssen und dabei die Energie drastisch minimieren. Würden beispielsweise 100 Teilnehmer, angeordnet in einem 10x10 Gitter, nicht mit der Zentralstation in der Mitte kommunizieren, sondern immer nur mit dem Partner in den vier Himmelsrichtungen erhielte man eine Energieeinsparung um den Faktor 9 (wenn zu den Gitternachbarn in gleichen Zeitabständen die gleiche Nachrichtenmenge fließt wie zuvor zur Zentralstation in der Mitte). Der Nachrichtenaufwand zur Übermittlung einer Nachricht würde bei 10 kurzen Übermittlungen statt einer 10-mal so langen Übertragung 10-mal weniger Energie benötigen.

Natürlich würde die Netzverfügbarkeit auch deutlich erhöht werden. Ein weiterer interessanter Vorteil von Ad-hoc-Netzwerken ist ihre Skalierbarkeit. Darunter versteht man die Fähigkeit eines Netzwerks seine Ressourcen an die Zunahme der Teilnehmerzahl anzupassen. Je mehr Teilnehmer sich anmelden, desto mehr Zwischenstationen stehen zur Verfügung und desto robuster und stabiler wird das Netzwerk. Das heißt zwar nicht, dass so ein Netzwerk nicht zusammenbrechen kann, aber es wird deutlich später passieren als in einem zellulären Netzwerk, wo die Kapazität durch die Platzierung und Ausstattung der Zentralstation vorgegeben ist.

Warum gibt es das noch nicht?

Tatsächlich sind die Hardwarevoraussetzungen von Mobil Ad-hoc-Netzwerken so gering, dass solche Netzwerke im Prinzip schon vor Jahrzehnten hätten realisiert werden können. Auch sind schon Algorithmen für solche Netzwerke länger bekannt. So kann das 1974 entwickelte Sequenced Distance Vector routing (DSDV) für Ad-hoc-Netzwerke angewendet werden. Und mittlerweile gibt es eine Reihe konkurrierender praxisreifer Konzepte für die Organisation und Unterhalt dieser Netzwerke. Hier eine kleine Auswahl:

- Destination Sequenced Distance Vektor (DSDV).

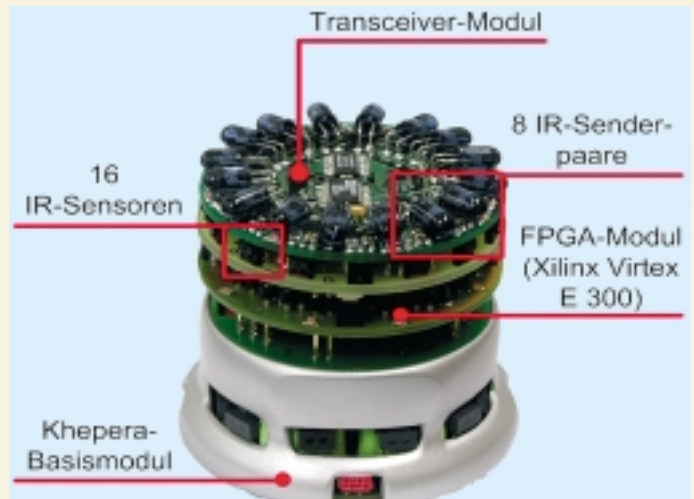


Abb. 4: Der Mini-Roboter Khepera und die Erweiterungsmodule mit Infrarot-Sende/Empfangskranz für den Einsatz in einem Mobilem Ad-hoc-Netzwerk.

- Optimized Link State Routing (OLSR).
- Dynamic Source Routing (DSR).
- Ad hoc On-demand Distance Vector (AODV).
- Temporally Ordered Routing Algorithm (TORA).
- Zone Routing Protocol (ZRP).
- Greedy Perimeter Stateless Routing (GPSR).

Der interessierte Leser sei für eine detaillierte Beschreibung dieser Protokolle auf das Buch von Perkins [6] verwiesen. All diese Ansätze funktionieren nur im kleinen Maßstab, d. h. bis zu einer Teilnehmerzahl von einigen hundert. Würde man diese Protokolle mit tausenden von Teilnehmern betreiben, so käme es zu mindestens einem der folgenden Probleme:

- Die Anzahl der Kontrollnachrichten nehmen überhand.
- Der Speicherverbrauch der Teilnehmerrechner wird zu groß.
- Netzteilnehmer werden nicht mehr gefunden.
- Kommunikationsverbindungen werden instabil.

Der Grund hierfür liegt in der mangelnden Fähigkeit der Netzwerke, skalierbare Lösungen anzubieten.

Hierzu ein Beispiel: In DSR werden Teilnehmer gesucht, indem an alle Teilnehmer Suchnachrichten versendet werden. Diesen Vorgang nennt man auch Fluten, weil quasi eine Woge von Nachrichten durch das Netzwerk schwappt. Sind nur wenige Teilnehmer im Netzwerk, so stellt dies kein Problem dar. Nur stelle man sich vor, dass in einem Netzwerk von 100 000 Teilnehmern jeder Teilnehmer bei der Herstellung jeder Verbindung assistieren muss. In anderen Lösungen wie DSDV werden riesige Tabellen über die momentanen Verbindungsnetze ausgetauscht. Man kann sich vorstellen, dass für kleine Netzwerke solche Netzwerkalgorithmen effizient sind. Für große Netzwerke ist der Informationsfluss nicht tragbar.

Peer-to-Peer-Netzwerke als Inspiration für das bessere Mobile Ad-hoc-Netzwerk

Man kann sich Peer-to-Peer-Netzwerke als das drahtgebundene Äquivalent von mobilen Ad-hoc-Netzwerken vorstellen. Die ersten populären Netze wie Napster und Gnutella skalierten jedoch ähnlich schlecht oder noch schlechter als die eben beschriebenen MANET-Protokolle. Dann kamen aber mit CAN, Chord und Tapestry moderne Peer-to-Peer-Netzwerk-Algorithmen auf, die beliebig große Netzwerke erlauben. Der entscheidende Trick war die Verwendung von verteilten Hash-Tabellen

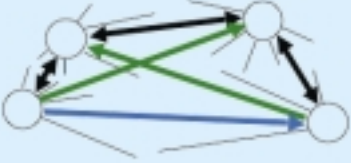
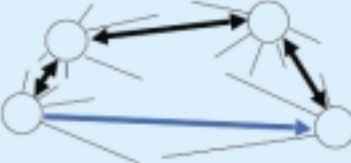
Netzwerktyp	Yao-Graph	Sparsified Yao-Graph (SparsY)	Symmetrischer Yao-Graph (SymmY)
Konstruktion	Wähle in jedem Sektor nächsten Nachbarn als Gesprächspartner	Nur die kürzeste eingehende Yao-Kante wird akzeptiert	Wähle aus Yao-Graph alle symmetrischen Kanten
Energieverbrauch	Niedrig	Niedrig	Niedrig
Congestion = Staugefahr	Hoch	Niedrig	Hoch
Grad = Verbindungen pro Sektor	Hoch	Zwei (gering)	Eins (minimal)
Interferenzen	Viele Interferenzen	Kaum Interferenzen	Keine Interferenzen
Beispiel			

Abb. 5: Die sektorbasierten Netzwerktopologie-Varianten des Yao-Graphs im Vergleich – eingesetzt wurde der SparsY-Graph.

(Distributed Hash Tables, DHT). Diese effiziente Datenstruktur erlaubt es jedem Teilnehmer, einen anderen kanonischen Teilnehmer zu finden, der diese Information speichert.

Hierzu ein Beispiel: Teilnehmerin Anna, die am 6. Juni 1967 geboren ist, ruft die Telefonnummer 060619670000 an und gibt dem verblüfften Boris, der das Gespräch annimmt, ihren Namen und ihre aktuelle Telefonnummer an. Will nun Christoph Anna anrufen (z. B. um zum Geburtstag zu gratulieren), dann ruft Christoph ebenfalls die 060619670000 an, um von Boris diese Information zu erhalten. Man beachte, dass sich so jeder Teilnehmer nur eine beschränkte Anzahl von Telefonnummern merken muss und dass die Last der Telefoneinträge (bei geeigneter Kodierung) gleichmäßig verteilt wird.

Skalierbare Mobile Ad-hoc-Netzwerke

Theoretische Untersuchungen des Autors an der Universität Paderborn zeigen, dass diese Technik der verteilten Hash-Tabellen sich auch auf mobile Ad-hoc-Netzwerke anwenden lässt. Ob sie sich auch praktisch umsetzen lässt, wird momentan durch die Implementation eines skalierbaren PaMANeT-Prototyps untersucht (Abbildung 3). Dies ist eine Kooperation im Institut für Informatik der Fachgebiete Algorithmen und Komplexität (Hochschuldozent Dr. Christian Schindelhauer, Dipl.-Inf. Gunnar Schomaker), Entwurf Paralleler Systeme (Prof. Dr. Franz-Josef Rammig, Dipl.-Inf. Simon Oberthür), Electronic-Commerce und Datenbanken (Prof. Dr. Stefan Böttcher, Dipl.-Inf. Adelhard Türling) im Rahmen einer Projektgruppe. Der Systementwurf des PaMANeT (Paderborn Mobile Ad-hoc-Network) baut auf den handelsüblichen W-LAN-Karten im IEEE 802.11-

Standard, wie man sie mittlerweile auch zu Hause einsetzt, auf und bietet eine Schnittstelle zu IPv6, der aktuellsten Version des Internet-Protokolls.

Wie gut kann solch ein MANET werden

Es liegt in der Natur des Menschen, nach den Grenzen einer neuen Technik zu fragen. Genau diese Frage wurde im Rahmen des Sonderforschungsbereichs SFB 376 Massive Parallelität der Universität Paderborn, Teilprojekt C6 Mobile Ad-hoc-Netzwerke untersucht. Die beiden Parameter Zeit und Energie interessieren zuerst.

Untersucht man die zeitbestimmenden Faktoren näher, entdeckt man zwei Ursachen für lange Kommunikationszeiten:

- Lange Wege, gemessen durch die Sprünge, die eine Nachricht vom Start zum Ziel zurücklegen muss und
- Staus, erzeugt durch störende andere Nachrichten, die in Routerknoten zu Verzögerungen oder im Äther zu Funkstörungen führen können.

Diese Ursachen bezeichnen wir mit Dilation und Congestion. Kann man aber für einen bestimmten Kommunikationsbedarf diese Übertragungen zeitgerecht durchführen? Entscheidend ist die Wahl der Kommunikationswege. Überträgt man zum Beispiel die Nachrichten immer direkt zum Ziel durch eine direkte Funkstrecke, so stören sich die Nachrichten gegenseitig durch das physikalische Phänomen der Funkinterferenz. Dadurch entstehen Rückstaus, da gewartet werden muss, bis die eine Frequenz frei wird. Wählt man aber Wege nur über die nächsten Nachbarn, so ergeben sich unnötig lange Wege und die Dilation wird entsprechend groß.

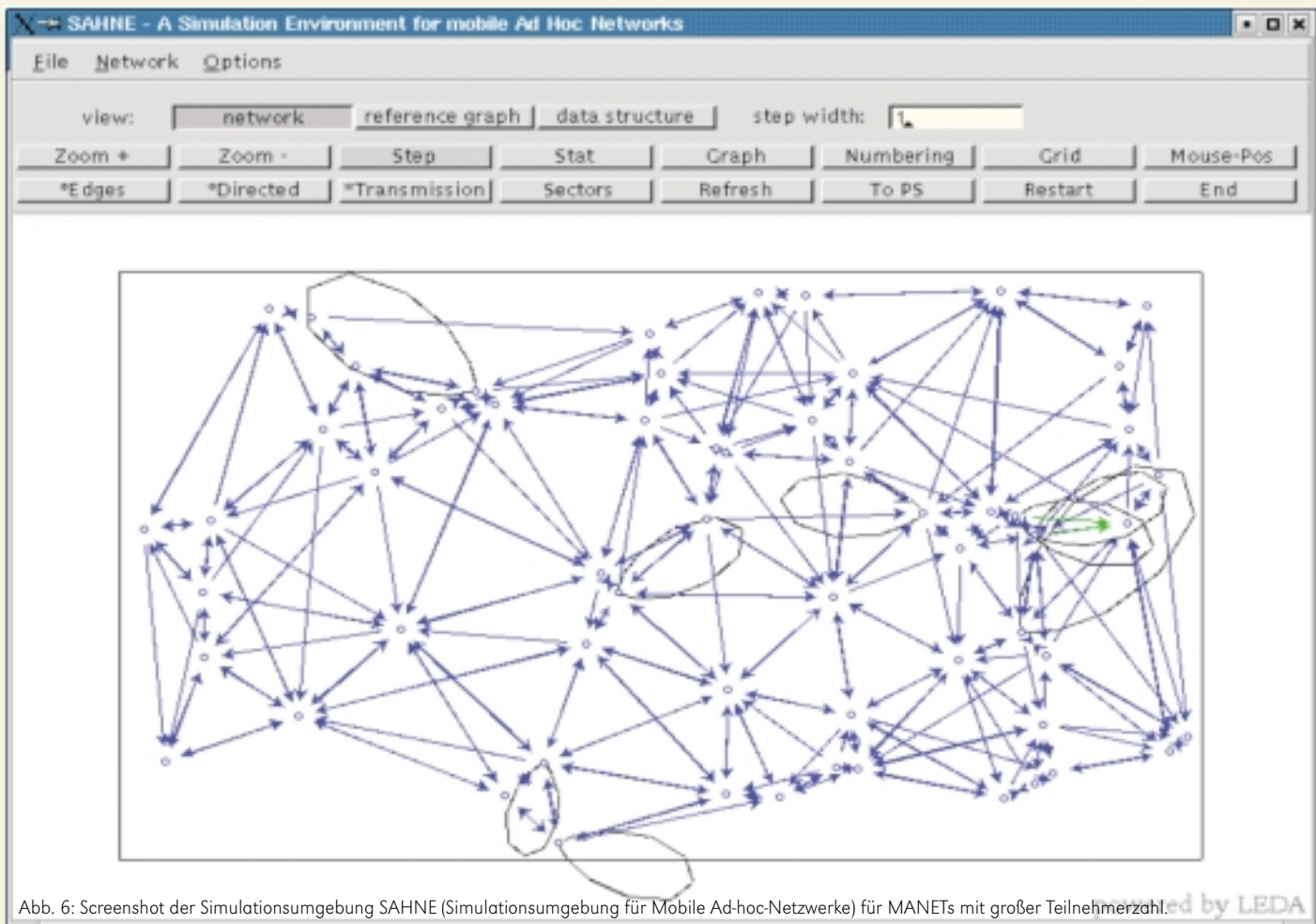


Abb. 6: Screenshot der Simulationsumgebung SAHNE (Simulationsumgebung für Mobile Ad-hoc-Netzwerke) für MANETs mit großer Teilnehmerzahl.

Theoretische Untersuchungen in [2] zeigen nun, dass zwischen Dilation und Congestion ein Trade-Off besteht. D. h. es gibt Situationen, in denen beide Maße nicht gleichzeitig optimiert werden können. Minimiert man die Dilation, so erhöht sich die Congestion und umgekehrt. Genauso ein Trade-Off ergibt sich für Energie und Dilation, was nicht sonderlich überrascht, da kurze Sprünge lange Funkstrecken notwendig machen mit entsprechendem zusätzlichen Energieverbrauch.

Vergleicht man Energie und Congestion, ist die Situation sogar verschärft. Hier kann nicht mehr von einem stufenlosen Übergang zwischen Energieoptimierung und Congestion-Minimierung gesprochen werden, die beiden Ressourcen sind unvereinbar. Es gibt keinen Kompromiss in der Optimierung dieser Maße.

Mit welchen Nachbarn soll man reden

Bevor man aber die Nachrichten durch geeignete Routenwahl verteilen kann, ergibt sich das Problem mögliche Kommunikationspartner auszuwählen. Zumeist wird dies durch Einschränkungen in der Kommunikationsausstattung notwendig. Bestehen zu viele Nachbarschaftsverbindungen, entsteht eine unnötige Menge von Kontrollnachrichten, um diese Verbindungen zu unterhalten. Werden bei dieser Linkauswahl Fehler gemacht, kann Folgendes passieren.

- Das Netzwerk zerfällt in nicht zusammenhängende Teile.
- Die Anzahl der Nachbarn eines Knotens ist zu groß.
- Unnötig lange Wege werden aufrechterhalten (Energie).

- Die Kommunikationsverbindungen stören sich gegenseitig zu stark (Congestion).
- Verbindungen mit wenigen Relaisstationen werden nicht unterstützt (Dilation).

Überraschenderweise kann man diese Nachbarschaftswahl mit einem einfachen Algorithmus zur Erzeugung des hierarchischen Schichtengraphs lösen und so ein Basisnetzwerk zur Verfügung stellen, in dem (fast) optimale Kommunikationsverbindungen für Energie, Congestion und Dilation hergestellt werden können. Nur kann man in diesem Netzwerk nicht mehr als einen dieser drei Parameter gleichzeitig optimieren (siehe Abbildung 5).

Infrarotkommunikation zwischen Mini-Robotern

Ein interessanter Spezialfall eines Mobilen Ad-hoc-Netzwerks ergibt sich, wenn die Teilnehmer Mini-Roboter sind, die sich über Infrarot-Verbindungen unterhalten (Abbildung 1 und 4). Dieses Szenario wurde im SFB 376 Massive Parallelität im Teilprojekt C6 untersucht in Zusammenarbeit der Fachgebiete Algorithmen und Komplexität mit Prof. Dr. Friedhelm Meyer auf der Heide und des Fachgebiets Schaltungstechnik des Instituts für Elektrotechnik mit Prof. Dr. Ulrich Rückert und Dipl.-Ing. Matthias Grünewald.

Hierbei unterscheidet sich die Infrarot-Kommunikation besonders in der Richtcharakteristik. Die Dioden strahlen das Licht in einem gewissen Sektor aus, so wie entsprechende Sensoren das Infrarotlicht nur in einem Sektor empfangen können. Um den gesamten Umkreis abzudecken, wurde daher auf die Miniroboter Khepera ein eigens konstruierter Sende- und Empfangskranz

aufgesetzt, so dass in jeder Richtung unabhängig kommuniziert werden kann. Wählt nun ein Roboter als Kommunikationspartner in jeder Richtung den am nächstgelegenen Nachbarroboter, erhält man als Kommunikationsnetzwerk den so genannten Yao-Graph.

Der Nachteil dieser nahe liegenden Struktur ist, dass ein Roboter potentiell der Ansprechpartner einer ganzen Reihe von Robotern sein kann, d. h. der Grad des Graphen ist zu groß. Daher wurde in [3] ein Alternativ-Konzept, der Sparsy-Graph (Sparsified Yao-Graph) vorgestellt. Hier akzeptiert jeder Roboter in einem Sektor nur eine eingehende Verbindung und zwar die des nächst gelegenen. Man kann nachweisen, dass so ein Basisnetzwerk sowohl hinsichtlich der Energie als auch der Congestion optimierte Kommunikation unterstützt (Abbildung 5). Dieses Konzept wurde daraufhin erfolgreich umgesetzt.

Weitere

aktuelle Fragestellungen

Wir haben gesehen, dass es im Bereich der Funknetzwerke durch die Ad-hoc-Netzwerke ein noch ungenutztes Potenzial gibt, das durch den Einsatz geeigneter Algorithmen bald genutzt werden kann. Es bleiben noch eine Reihe offener Forschungsfragen, von denen einige nicht unerwähnt bleiben sollen.

- **Routing:**

Das Thema Routenwahl für Ad-hoc-Netzwerke ist noch nicht abgeschlossen und große Hoffnungen bestehen insbesondere im Zusammenhang mit den aus Paderborn stammenden preisgekrönten Routingverfahren von Harald Räcké [4].

- **Testen großer Ad-hoc-Netzwerke:**

Man weiß wenig über das tatsächliche Verhalten von Ad-hoc-Netzwerken mit großer Teilnehmerzahl. Mit der Entwicklung der Simulationsumgebung (SAHNE – Simulation Environment for Mobile Ad hoc Networks) [5] gelang es zwar weitaus größere Netzwerke zu simulieren als jemals zuvor. Simulationen werden aber einen Praxistest nicht ersetzen können (Abbildung 6).

- **Mobilität:**

In [1] wurden die theoretischen Grundlagen für eine praxisnahe Modellierung von Mobilität in Ad-hoc-Netzwerken gelegt. Das Verständnis der Mobilität in MANETs steckt aber immer noch in den Kinderschuhen.

Literatur

Tamás Lukovszki, Stefan Rührup, Christian Schindelhauer, Klaus Volbert: Worst Case Mobility in Ad-Hoc Networks, ACM Symposium on Parallelism in Algorithms and Architectures, 2003 (SPAA 2003), S. 230-239.

Friedhelm Meyer auf der Heide, Christian Schindelhauer, Klaus Volbert, Matthias Grünewald: Congestion, Energy and Delay in Radio Networks, Fourteenth ACM Symposium on Parallel Algorithms and Architectures SPAA, 2002, 230-237.

Matthias Grünewald, Tamás Lukovszki, Christian Schindelhauer, Klaus Volbert: Distributed Maintenance of Resource Efficient Wireless Network Topologies, 8th International Euro-Par Conference Paderborn, Germany, August 2002 (Euro-Par 2002 Parallel Processing), 935-946, (distinguished paper).

Harald Räcké: Minimizing Congestion in General Network, 43th Annual IEEE Symposium on Foundations of Computer science, Vancouver, Canada, November 2002, S. 43-52.

Klaus Volbert: A Simulation Environment for Ad Hoc Networks Using Sector Subdivision, in Proc. of the 10th Euromicro Workshop on Parallel, Distributed and Network-based Processing (PDP'02), 2002, S. 419-426.

Charles E. Perkins: Ad Hoc Networking. Addison-Wesley, 2001.

DEN KOPF VOLLER IDEEN, EIN KLARES ZIEL VOR AUGEN.



**RIGHT
CHAIRING**

Wer kann, der darf: Ob Sie im Rahmen eines Praktikums erstmals Berufsalltag schnuppern, Ihre Diplomarbeit bei uns schreiben oder mit abgeschlossener Ausbildung bei uns starten wollen – wir haben für ehrgeizige Einsteiger immer einen Stuhl frei. Als international ausgerichtetes Unternehmen suchen wir Denker, Macher und Talente aus unterschiedlichen Bereichen. Und weil wir 75% unserer Führungskräfte aus den eigenen Reihen besetzen wollen, stehen Ihre Chancen bei uns auch langfristig gut.

Nachwuchskräfte für unterschiedliche Fachbereiche

- Praktikanten
- Diplomanden
- Absolventen

Weidmüller ist der führende Hersteller von Komponenten für die elektrische Verbindungstechnik. Zu dem Weidmüller-Produktportfolio zählen Reihenklemmen, Steck- und Leiterplattenverbinder, geschützte Baugruppen sowie Relaiskoppler bis hin zu Stromversorgungs- und Überspannungsschutz-Modulen in allen Anschlussarten. Material zur Elektroinstallation und Betriebsmittelkennzeichnung, E/A-Basiskomponenten und Werkzeuge runden das Programm ab. Als OEM-Anbieter setzt das Unternehmen dabei weltweit Standards in der elektrischen Anschluss- und Verbindungstechnik. Weltweit beschäftigt Weidmüller derzeit insgesamt rund 2.300 Mitarbeiter und ist in mehr als 70 Ländern für seine Kunden tätig. Weidmüller erzielte im Geschäftsjahr 2003 einen Umsatz von 326 Mio. Euro.

Weidmüller Interface GmbH & Co. KG – Akademie

Schul- und Hochschulbetreuung

Postfach 30 30, 32760 Detmold

Bewerberhotline: 0 52 31 / 14 - 18 74

E-Mail: hochschulbetreuung@weidmueller-akademie.de

Gehen Sie uns ins Netz: www.weidmueller.com

Wer alles gibt, gibt nie zu wenig

Weidmüller 

ICE-Radreifenbruch

Bruchmechanik trägt zur Klärung der Schadensursache bei

Am 3. Juni 1998 verunglückte in Eschede der ICE Wilhelm-Conrad Röntgen auf seiner Fahrt von München nach Hamburg mit den allseits bekannten schwerwiegenden Folgen. Mehrere Wagen des Zuges waren mit gummigefeder-ten Rädern ausgestattet. Eines dieser Räder zerbrach ca. 6 km vor einer Weiche. An dieser Weiche stellten sich Teile des Zuges quer und prallten gegen die nachfolgende Brücke, die dabei einstürzte. Als Unfallursache wurde der gebrochene Radreifen des gummigefederten Rades ausgemacht. Gummigefederte Räder bestehen im Wesentlichen aus dem Radreifen, der Felge und den dazwischen verspannten Gummikörpern. Die wesentliche Kraftübertragung findet zwischen Radaufstandspunkt und Radwelle statt, wobei die höchste Spannung im Radreifen auf der Innenseite des Radreifens auftritt. Infolge der Radumdrehungen kommt es zu einer zyklischen Beanspruchung, die zu einem ausge-dehten Ermüdungsrisswachstum und letztlich zum Bruch des Radreifens führte.

Im Folgenden wird insbesondere auf numerische und experimen-telle Simulationen des Ermüdungsrisswachstums eingegangen, die im Zusammenhang mit den Schadensuntersuchungen in der Fachgruppe Angewandte Mechanik der Universität Paderborn in Kooperation mit dem Westfälischen Umweltzentrum durchge-führt wurden.

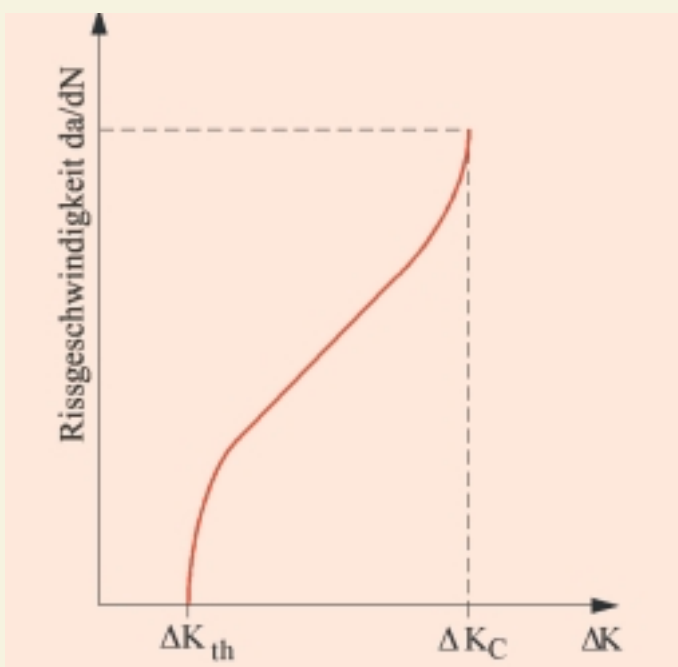


Abb. 1: Zusammenhang zwischen Rissgeschwindigkeit da/dN und zyklischem Spannungsintensitätsfaktor ΔK .



Prof. Dr.-Ing. Hans Albert Richard ist Professor für Angewandte Mechanik an der Universität Paderborn. Er promovierte 1979 und habilitierte 1984 an der Universität Kaiserslautern. 1986 folgte er dem Ruf an die Universität Paderborn, an der er von 1991 bis 1995 Rektor war. Prof. Richard ist Vorsitzen-der des Westfälischen Umwelt Zentrums (WUZ). Im Jahr 2002 wurde er vom Landge-richt Lüneburg als Gutachter im Prozess um das ICE-Unglück von Eschede bestellt. 2004 erhielt er mit der GRIFFITH-Medaille die welt-weit höchste Auszeichnung auf dem Gebiet der Bruchmechanik.

Grundlagen

des Ermüdungsrisswachstums

Bruchmechanische Untersuchungen gehen von der Existenz von Rissen in Bauteilen und Strukturen aus. Diese Risse stören das vorliegende Spannungsfeld und vermindern die Tragfähigkeit und die Lebensdauer von Konstruktionen. Die Gefährlichkeit eines Risses kann beurteilt werden durch den Spannungsinten-sitätsfaktor K . Bei zeitlich veränderlicher Belastung ist der so genannte zyklische Spannungsintensitätsfaktor ΔK entscheidend. Dieser lässt sich nach der Formel

$$\Delta K = \Delta \sigma \sqrt{\pi a Y} \quad (1)$$

aus der zyklischen Spannung $\Delta \sigma$, der Risslänge a und dem Geometriefaktor Y errechnen [1].

Bei Ermüdungsbelastungen stellen sich u. a. folgende Fragen:

- Ist der Riss ausbreitungsfähig?
- Mit welcher Geschwindigkeit breitet sich der Riss aus?
- Bei welcher Risslänge bzw. welcher Belastung wird der Riss instabil?
- Wie groß ist die Lebensdauer bzw. die Restlebensdauer des Bauteils?

Diese Fragen lassen sich mit der von Werkstoff zu Werkstoff unterschiedlichen da/dN - ΔK -Kurve (Abbildung 1) beantworten, die den Zusammenhang zwischen der Rissgeschwindigkeit da/dN (Risslängenänderung pro Zyklenzahl der schwingenden Belas-tung) und dem zyklischen Spannungsintensitätsfaktor ΔK , Gleichung (1), darstellt. Ermüdungsrissausbreitung ist demnach möglich, wenn die zyklische Spannungsintensität ΔK größer als ΔK_{th} und kleiner als ΔK_C ist:

$$\Delta K_{th} < \Delta K < \Delta K_C \quad (2).$$

ΔK_{th} ist dabei der Schwellenwert der Ermüdungsrissausbreitung. Für $\Delta K < \Delta K_{th}$ ist ein Riss nicht ausbreitungsfähig. ΔK_C stellt die zyklische Spannungsintensität dar, bei der Bruch, d. h. instabiles Risswachstum eintritt. Die sich für die aktuell vorliegende Rissbe-



Abb. 2: Unterschied zwischen Ermüdungsbruch- und Restgewaltbruchfläche bei einem hochfesten Stahl.

lastung ΔK ergebende Rissgeschwindigkeit kann unmittelbar aus dem Diagramm in Abbildung 1 abgelesen werden. Die Restlebensdauer eines Bauteils erhält man durch Integration der dargestellten $da/dN-\Delta K$ -Kurve.

Ermüdungsrisswachstum unterscheidet sich wesentlich bei zyklischer Belastung mit konstanter Amplitude und zyklischer Belastung mit variabler Amplitude. Bei konstanter Amplitude ergibt sich ein kontinuierliches Risswachstum, während die bei variabler Amplitude auftretenden Reihenfolgeeffekte für ein sehr diskontinuierliches Risswachstum sorgen [2].

Bei Schadensfällen geben die Bruchflächen Aufschluss über das erfolgte Risswachstum und die tatsächliche Belastung einer Konstruktion. So lässt sich das Gebiet des stabil verlaufenden Ermüdungsrisswachstums sehr deutlich von dem instabilen Restgewaltbruch unterscheiden (Abbildung 2). Ist die Ermüdungsbruchfläche klein und die Restgewaltbruchfläche groß, so ist von hoher Bauteilbelastung auszugehen und umgekehrt.

Auf der Ermüdungsbruchfläche ist auch zu erkennen, ob es sich um eine Belastung mit konstanter oder mit variabler Amplitude handelt.

Aufbau, Belastung und Tragverhalten gummigefederter Räder

Gummigefederte Räder werden seit Jahrzehnten weltweit bei Fahrzeugen des Straßen-, Stadt- und U-Bahnverkehrs verwendet, um z. B. Verschleiß und Geräuschentwicklungen, bedingt durch enge Gleisbögen und andere Besonderheiten dieser Schienennetze, in Grenzen zu halten.

Aufgrund von Schwingungen und Geräuschen in ICE-Zügen wurden gummigefederte Räder auch für den Vollbahnbereich weiterentwickelt und im Herbst 1992 für den Hochgeschwindigkeitsverkehr der Deutschen Bahn zugelassen. Diese neue Radsatzbauart mit der Bezeichnung BA 064 hat im Neuzustand einen Messkreisdurchmesser von 920 mm. Durch Verschleiß und Reprofilierung verringert sich der Durchmesser erheblich. Das Betriebsgrenzmaß wurde bei einem Raddurchmesser von 854 mm festgesetzt. Bei dem Unfallrad lag der Messkreisdurchmesser bei 862 mm, d. h. die Dicke des Radreifens hatte sich von 60 mm auf 31 mm reduziert.

Ein gummigefedertes Rad, wie es beim ICE zum Einsatz kam, besteht aus einem Radreifen, 34 Gummikörpern, einem Radkörper (Felge) und einer Vollwelle (Abbildung 3). Zum Zwecke der Montage ist die Felge zweigeteilt, sie besteht aus einem Felgenkörper und einem Felgenring. Zunächst werden die Gummikörper in gleichmäßigem Abstand in den Spalt zwischen Radreifen und Radkörper eingebracht. Zuletzt wird der Felgenring mit dem Radkörper verschraubt. Dabei werden die Gummikörper verspannt, d. h. sie werden in radialer und in axialer Richtung des Rades gestaucht und dehnen sich in Umfangsrichtung in die bestehenden Freiräume aus. Durch die Reibung zwischen den Gummikörpern und dem Radreifen bzw. der Felge sind die

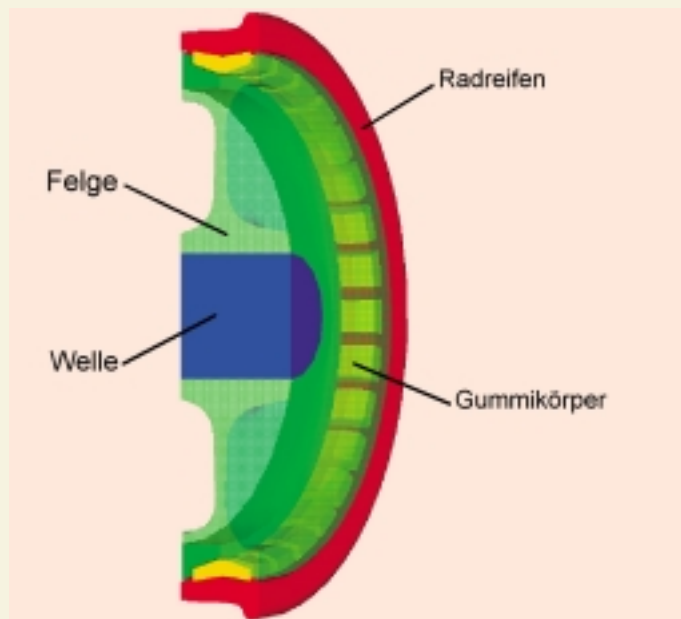


Abb. 3: Aufbau eines gummigefederter Rades (CAD-Modell einer Radhälfte).

Gummikörper gegen Verrutschen gesichert.

Das montierte, gummigefederte Rad stellt somit eine Gesamttragstruktur dar. Die wesentliche Kraftübertragung findet zwischen Radaufstandspunkt und Radsatzwelle statt. Dabei wird die Radaufstandskraft über den Radreifen auf die Gummielemente und von dort in die Radfelge und somit auf die Radwelle übertragen (Abbildung 4).

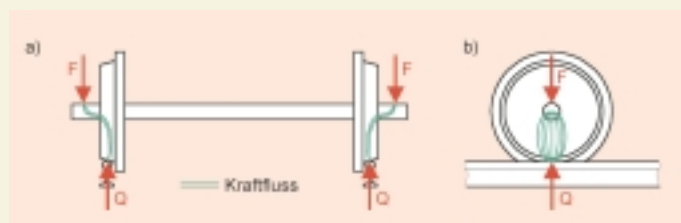


Abb. 4: Radsatz mit Belastung und prinzipiellem Kraftflussverlauf beim Lastfall Geradeausfahrt – a) Vorderansicht, b) Seitenansicht.

Nach dem UIC-Entwurf 510-5 [3] bzw. DIN EN 13979-1 [4] unterscheidet man insbesondere die Lastfälle „Geradeausfahrt“, „Kurvenfahrt (Bogenfahrt)“ und „Weichenfahrt“. Für die Geradeausfahrt ergibt sich die Radaufstandskraft

$$Q = 1,25 Q_0 = 98 \text{ kN.}$$

Q_0 stellt die statische Radlast dar, die sich aus dem Wagengewicht ergibt.

Die für die Rissausbreitung entscheidende höchste Umfangsspannung tritt am Innenrand des abgefahrenen Radreifens, d. h. bei minimalem Messkreisdurchmesser, beim Lastfall Geradeausfahrt auf.

Rechnerische Spannungsanalyse

Für eine genaue rechnerische Spannungsanalyse von Radreifen ist eine dreidimensionale Finite-Elemente-Analyse erforderlich. Dies ist bedingt durch die Geometrie der Räder, die mehrdimensionale Verspannung der Gummikörper und die räumliche Belastungssituation, insbesondere, wenn man alle beim Schienenverkehr auftretenden Lastfälle betrachtet.

Die nachfolgenden Untersuchungen beziehen sich auf die Abmessungen des Unfallrads mit einem Durchmesser von 862 mm. Abbildung 5 zeigt das Finite-Elemente-Netz für das

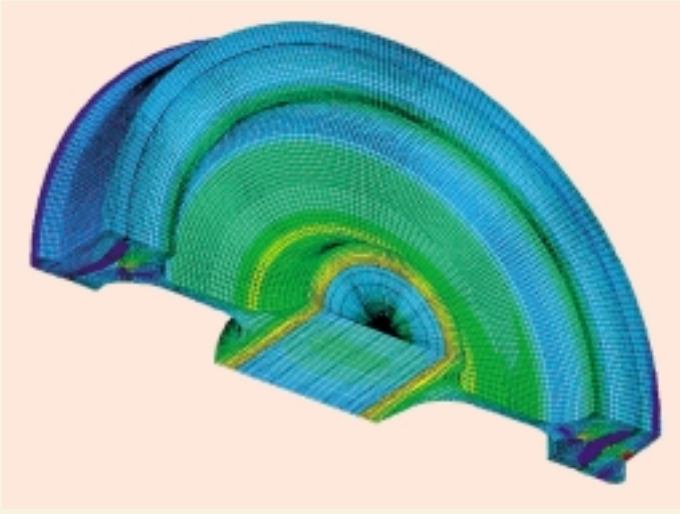


Abb. 5: Finite-Elemente-Netz für das ICE-Rad mit einem Durchmesser von 862.

ICE-Rad. Aus Symmetriegründen genügt es nur das halbe Rad zu betrachten. Die Symmetriebedingungen in der Schnittfläche und die Lagerungen werden durch entsprechende kinematische Randbedingungen erfüllt.

Das in Abbildung 5 dargestellte Finite-Elemente-Netz besteht aus ca. 130 000 Elementen mit mehr als 150 000 Knoten. Bei dem System „Gummigefedertes Rad“ liegt eine dreifache Nichtlinearität vor:

- Nichtlinearität aufgrund des nichtlinearen Werkstoffverhaltens der Gummikörper.
- Geometrische Nichtlinearität aufgrund der großen Verformungen der Gummikörper.
- Strukturnichtlinearität aufgrund des Kontakts zwischen den Radteilen.

Bei der Finite-Elemente-Analyse ist ein zweigestuftes Vorgehen erforderlich. Zunächst müssen das Rad rechnerisch montiert und die daraus entstehenden Montagespannungen ermittelt werden. Dann erfolgt die Spannungsanalyse infolge der Radbelastung. Aufgrund der Nichtlinearität des Systems „Gummigefedertes Rad“ muss das Gleichungssystem mit mehr als 450 000 Gleichungen für die Montage ca. 60-mal und für jeden Lastfall ca. 50-mal gelöst werden.

Durch die Montage tritt eine Umfangsspannung auf, die in Umfangsrichtung, abgesehen von kleinen Schwankungen durch die Gummielemente, konstant und in axialer Richtung aufgrund der unsymmetrischen Form des Radreifenquerschnitts in der Höhe veränderlich ist. An der Innenseite des Radreifens beträgt die Umfangsspannung infolge der Montage etwa 50 MPa.

Die Analyse des Lastfalls „Geradeausfahrt“ ergibt am Innenrand des Radreifens positive Umfangsspannungen, die sich in Umfangsrichtung und axialer Richtung sehr stark ändern. Die höchste Umfangsspannung tritt an der Innenseite des Radreifens in der Symmetrieebene des belasteten Rades auf.

Für den Lastfall Geradeausfahrt ergibt sich für die Radaufstandskraft $Q = 98 \text{ kN}$ eine maximale Spannung $\sigma_{\max}$ von 220 MPa und eine minimale Spannung $\sigma_{\min}$ von 6 MPa.

Abbildung 6 zeigt die Umfangsspannung auf dem Innenrand des Radreifens bei 2 Radumdrehungen. Man erkennt, dass bei jeder Radumdrehung ein Belastungszyklus durchlaufen wird. Für eine Radlast von $Q = 98 \text{ kN}$ ergibt sich somit eine Spannungsamplitude von $\sigma_a = 107 \text{ MPa}$, eine Mittelspannung von $\sigma_m = 113 \text{ MPa}$ und ein R-Verhältnis von $R = +0,03$.

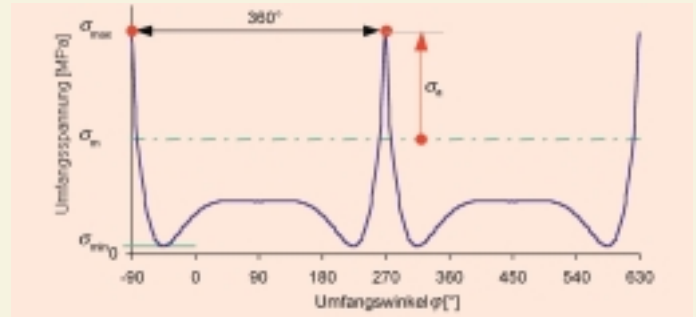


Abb. 6: Umfangsspannung auf dem Innenrand des Radreifens bei zwei Radumdrehungen.

Schadensanalyse des Radreifenbruchs

Bei zyklischer Belastung (Ermüdungsbelastung) von Bauteilen und Strukturen setzt sich die Gesamtlebensdauer aus einer Periode der Risserzeugung (Rissbildung) und einer Periode des Risswachstums zusammen. Bei einer gleichmäßig belasteten Struktur mit glatter Oberfläche liegen eine sehr lange Rissentstehungsphase und eine relativ kurze Risswachstumsphase vor. Bei Bauteilen mit ungleichmäßiger Spannungsverteilung an der Oberfläche und vorhandenen Oberflächen- oder Materialfehlern nimmt dagegen das Risswachstum eine längere Periode der dann niedrigeren Gesamtlebensdauer ein.

Da beim gummigefederten Rad der innere Rand des Radreifens in Umfangsrichtung ohne Kerben ist, erfährt jeder Radpunkt bei konstanter Radlast Q eine gleich hohe zyklische Belastung, die sich bei jeder Radumdrehung wiederholt (Abbildung 6). Systematische Untersuchungen am gebrochenen Radreifen haben ergeben, dass keine Oberflächenfehler und Materialfehler am Innenrand des Radreifens vorlagen. Es kann somit von einer langen Risserzeugungsphase ausgegangen werden. Dies bedeutet nicht, dass die Risswachstumsphase kurz, sondern nur, dass die Risswachstumsphase voraussichtlich kürzer als die Rissentstehungsphase war.

Das Risswachstum beim ICE-Rad begann in der Nähe der bei dem Lastfall Geradeausfahrt auftretenden maximalen Spannung am Innenrand des Radreifens (Abbildung 7). Der Riss wuchs zunächst mehr in die Tiefe des Radreifens, um sich später in halbelliptischer Form auszubreiten. Dabei wuchs der Riss dann erheblich schneller in die Breite als in die Tiefe. Erst als der Reifenquerschnitt zu ca. 80 Prozent durch das Risswachstum geschädigt war, trat der Bruch des Radreifens ein.

Auf der Bruchfläche erkennt man Rastmarken, Farbeffekte und sonstige Bruchflächenstrukturen, die auf sehr diskontinuierliches Risswachstum hindeuten. So wechseln sich Phasen sehr langsamen Risswachstums oder Rissstillstands und Phasen schnelleren Risswachstums ab. Insgesamt deutet alles auf ein zeitlich sehr ausgedehntes stabiles Risswachstum hin.

Das Bruchbild (Abbildung 7) zeigt sehr deutlich, dass der Werkstoff des Radreifens sehr risstolerant ist. In vielen Bauteilen und Strukturen setzt bereits bei relativ kleinen Ermüdungsrissen instabiles Risswachstum ein. Nicht so beim ICE-Rad, hier konnte der stabile Riss eine Querschnittsfläche von ca. 80 Prozent einnehmen, bevor Bruch eintrat. Diese Überlegungen werden auch gestützt durch die Tatsache, dass der Risszähigkeitswert für das Radreifenmaterial $K_{IC} = 86,8 \text{ MPa}\sqrt{\text{m}}$ beträgt und somit relativ hoch liegt. Anhand numerischer und experimenteller Simulationen wird das Ermüdungsrisswachstum im Radreifen im Nachfolgenden näher untersucht.

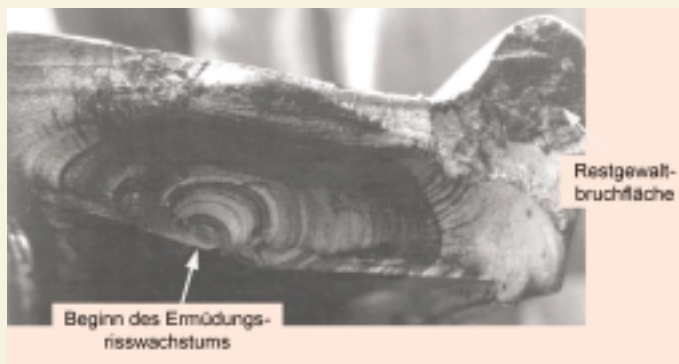


Abb. 7: Bruchfläche des gebrochenen Radreifens.

Numerische Simulation des Ermüdungsrissswachstums

Das Ermüdungsrissswachstum in dreidimensionalen Strukturen lässt sich mit der Finite-Elemente-Methode simulieren. Besonders bewährt hat sich hier das Programmsystem ADAPCRACK3D [5, 6], das in der Fachgruppe Angewandte Mechanik der Universität Paderborn entwickelt wurde. Das Programm ADAPCRACK3D (Abbildung 8) besteht aus 3 Funktionsmodulen, die gemeinsam eine automatische Risssfortschrittssimulation durchführen. Als Eingaben werden im Wesentlichen eine Beschreibung des ungerissenen Bauteils in Form eines 3D-Finite-Elemente-Netzes, eine Anrissbeschreibung sowie Materialdaten zur Definition des Rissausbreitungsverhaltens benötigt. Das Modul NETADAPT3D leistet alle notwendigen Netzanpassungen, die sich durch die Änderung der Bauteilgeometrie infolge der Rissausbreitung ergeben. In jedem Schritt wird eine Finite-Elemente-Berechnung durchgeführt. Das Modul NETCRACK3D berechnet aus der FE-Lösung die Spannungsintensitätsfaktoren an der Rissfront. Diese dienen einerseits zur Berechnung neuer Rissfrontkoordinaten und andererseits zur Bestimmung der für ein definiertes Rissinkrement benötigten Lastwechselzahl. Die automatische Rissswachstumssimulation wird solange fortgesetzt bis z. B. die Risszähigkeit des Materials erreicht wird und somit instabiles Rissswachstum eintritt.

Zur Simulation des Rissswachstums im ICE-Radreifen mit einem Raddurchmesser von 862 mm wurde ein Segment des ungerissenen Radreifens mit 52 000 Finite-Elementen diskretisiert. In dieses FE-Modell wurden in einem ersten Schritt Anrisse unterschiedlicher Länge und Tiefe eingebracht, um eine geeignete Anrissgröße für die Risssfortschrittssimulation zu bestimmen. Für die Anrissimulation wird dabei von einem halbkreisförmigen Anriss ausgegangen. Durch die Vorsimulationen wird ein Anriss mit einem Radius von $r = 1,5 \text{ mm}$ ermittelt, bei dem entlang der gesamten Rissfront die Spannungsintensitäten den Schwellenwert $\Delta K_{th} = 8,2 \text{ MPam}^{1/2}$ des Radreifenmaterials überschreiten. Da das Ermüdungsrissswachstum beim Radreifenbruch nicht am „Dachfirst“, sondern 13 mm versetzt begann, wird der 1,5 mm tiefe Anriss entsprechend angeordnet (Abbildung 9a). Die Simulation wird für eine Radlast von $Q = 98 \text{ kN}$ durchgeführt. Dabei ergibt sich entsprechend der zyklischen Belastung je Radumdrehung (Abbildung 6) ein gewisser Risssfortschritt. Die gesamte Risssfortschrittssimulation umfasst 26 Simulationsschritte und wird bis zum Erreichen der Risszähigkeit von $K_C = 86,8 \text{ MPam}^{1/2}$ durchgeführt.

Die sich je Simulationsschritt ergebenden Rissfronten sind in Abbildung 9b dargestellt. Der Riss wächst zunächst eher halb-

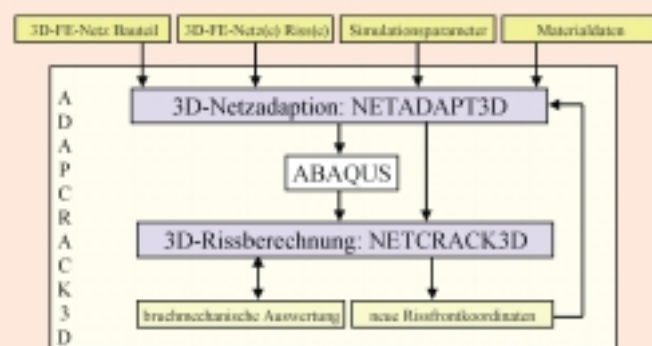


Abb. 8: Aufbau des Programmsystems ADAPCRACK3D.

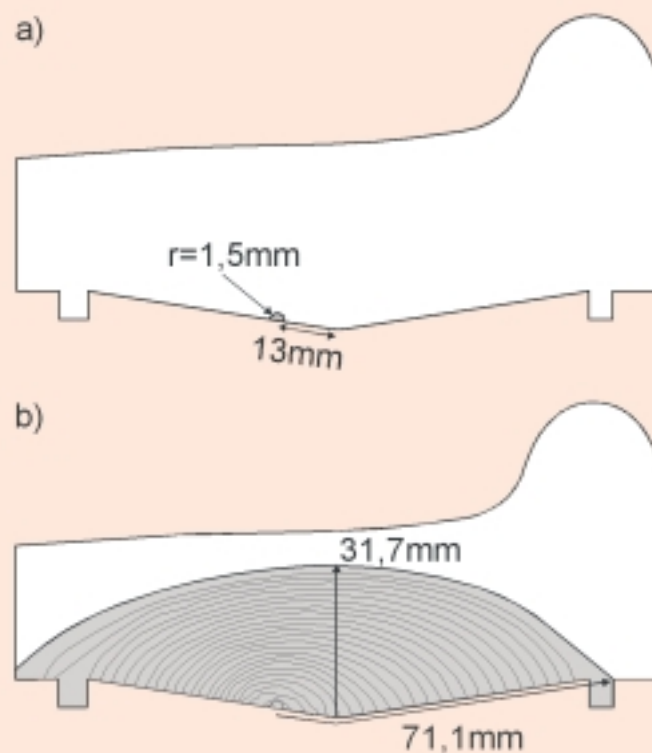


Abb. 9: Numerische Risssimulation für den ICE-Radreifen mit 862 mm Durchmesser (Unfallrad) und einer Q-Kraft von 98 kN – a) Anrisslage und Rissabmessung bei Simulationsbeginn, b) Simulierte Rissfronten mit den Rissabmessungen unmittelbar vor Eintritt der Instabilität.

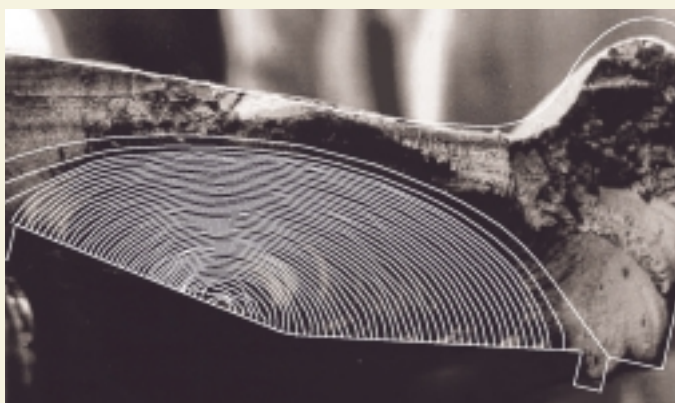


Abb. 10: Vergleich des simulierten Rissswachstums mit dem tatsächlichen Rissswachstum beim Radreifenbruch.

kreisförmig, um sich später deutlich schneller in die Breite auszuweiten. Unmittelbar vor Eintritt der Instabilität hat der Riss eine Tiefe von 31,7 mm und eine maximale Risslänge an der Radreifenninnenseite von 71,1 mm. Abbildung 10 zeigt einen Vergleich

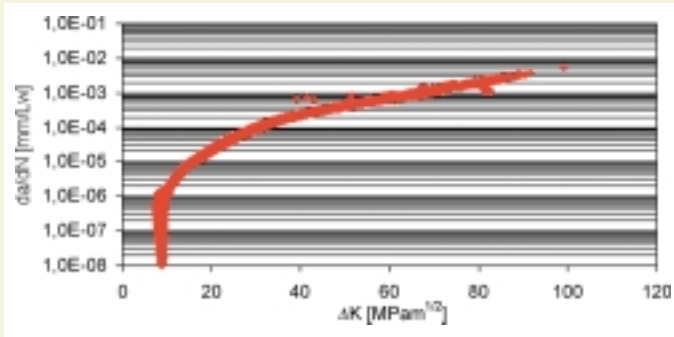


Abb. 11: Rissgeschwindigkeitskurve für den Radreifenwerkstoff (B5).

des simulierten Risswachstums mit dem tatsächlichen Risswachstum beim Radreifenbruch. Es wird deutlich, dass die numerische Simulation sehr gut geeignet ist, derartige Schadensereignisse vorherzusagen. Es bleibt festzuhalten, dass der Bruch des ICE-Radreifens erst nach erheblichem stabilem Risswachstum bei nur sehr kleiner Restgewaltbruchfläche erfolgt ist.

Die Risswachstumslebensdauer ergibt sich unter der Annahme einer Anrisslänge von 1,5 mm und einer zyklischen Belastung mit konstanter Amplitude und einer konstanten Maximallast von $Q_{\text{max}} = 98 \text{ kN}$ unter Berücksichtigung der da/dN -Kurve nach Abbildung 11 rein rechnerisch zu ca. 1,4 Mio. Lastwechseln. Dabei liegt eine lineare Schadensakkumulationsrechnung zugrunde, die keine Verzögerungen durch Reihenfolgeeffekte berücksichtigt. Die Bruchfläche (Abbildung 7) zeigt aber deutliche Spuren eines diskontinuierlichen Risswachstums. Die reale Risswachstumslebensdauer kann daher eher durch Experimente herausgefunden werden.

Experimentelle Simulation des Ermüdungsrisswachstums

Ermüdungsrisswachstum erfolgt bei zyklischer Belastung unter bestimmten Voraussetzungen. Charakteristisch für das Ermüdungsrisswachstum ist die Rissgeschwindigkeit da/dN und die zyklische Spannungsintensität ΔK . Dieser Zusammenhang wird in der Rissfortschrittskurve dargestellt und experimentell nach der amerikanischen Norm ASTM E 647 mittels geeigneter Proben, z. B. der CT-Probe, ermittelt. Abbildung 11 zeigt die Rissgeschwindigkeitskurve $da/dN = f(\Delta K)$ für den Radreifenwerkstoff. Als Schwellenwert der Ermüdungsrissausbreitung ergaben die Messungen $\Delta K_{\text{th}} = 8,2 \text{ MPa}\sqrt{\text{m}}$, für die Risszähigkeit wurde $K_{\text{IC}} = 86,8 \text{ MPa}\sqrt{\text{m}}$ ermittelt.

Bei der Fachgruppe Angewandte Mechanik (FAM) der Universität Paderborn werden die Versuche computergesteuert mit dem Programm FAMControl durchgeführt [7], das eigens dafür von der FAM entwickelt wurde. Zur Aufnahme einer kompletten da/dN - ΔK -Kurve sind mindestens 2 Versuche erforderlich [8]. Mit einem Versuch wird durch lineares Absenken der Spannungsintensität der mittlere und untere Teil der Rissfortschrittskurve ermittelt. In einem weiteren Versuch wird durch eine kontinuierliche Zunahme der Spannungsintensität der mittlere und obere Bereich der Kurve bestimmt.

Bei diesen Versuchen zeigt sich, dass die Ermüdungsbruchfläche bei höherer Rissgeschwindigkeit deutlich heller ist als bei niedriger Rissgeschwindigkeit. Auch ist ein deutlicher Unterschied zwischen relativ glatter Ermüdungsbruchfläche und sehr grober Restgewaltbruchfläche sichtbar.

Helle und dunkle Einfärbungen zeigen sich auch bei zyklischer

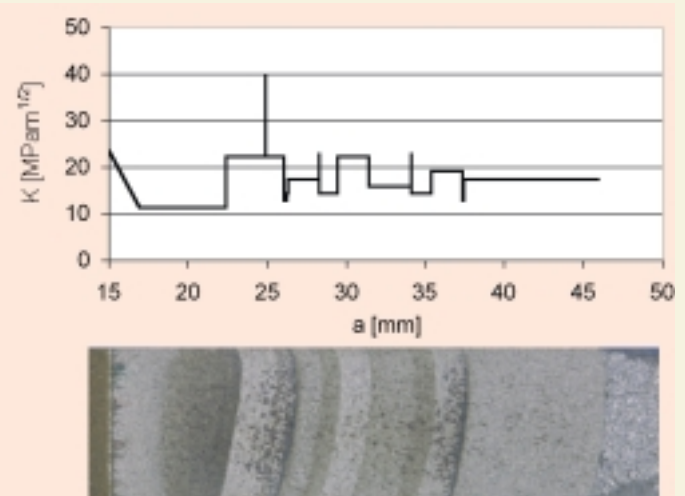


Abb. 12: Zusammenhang zwischen Belastung und Bruchfläche.

Belastung mit variabler Amplitude. So führen Änderungen in der Belastungshöhe zu Rastmarken oder helleren und dunkleren Oberflächenstrukturen (Abbildung 12).

Derartige Markierungen sind auch auf der Bruchfläche des ICE-Radreifens zu entdecken (Abbildung 7). Auf der Basis der Einfärbungen der Bruchfläche des Radreifens wurden anhand der Bruchflächen der CT-Proben, die für die Ermittlung der da/dN -Kurven zum Einsatz kamen, die Rissgeschwindigkeiten und die Spannungsintensitätsfaktoren zugeordnet. Diese Erkenntnisse sind in Nachfahrversuche mit CT-Proben eingeflossen. Für ein Risswachstum von 30 mm ergibt sich dabei eine Lebensdauer von mehr als 35 Mio. Lastwechsel. Dies entspricht einer Fahrstrecke von ca. 95 000 km für den ICE. Abbildung 13 stellt einen Vergleich der Bruchflächen für den Radreifen und für den Nachfahrversuch mit der CT-Probe dar. Die gute Übereinstimmung deutet daraufhin, dass beim Radreifenbruch ein sehr diskontinuierliches Risswachstum stattgefunden hat [9].



Abb. 13: Vergleich der Bruchflächen für den Radreifen und die Probe.

Schlussbemerkungen

Die Untersuchungen zeigen, dass mit der Bruchmechanik das Risswachstum im ICE-Radreifen gut beschrieben werden kann. Die in Paderborn durchgeführten Untersuchungen haben damit wesentlich zur Schadensaufklärung im ICE-Prozess vor dem Landgericht Lüneburg beigetragen.

Danksagung

Der Autor bedankt sich bei Frau Dr. M. Sander, Herrn G. Kullmer und Herrn M. Fulland für die wertvolle Unterstützung bei den Untersuchungen zum ICE Radreifenbruch.

Literatur

- [1] H. A. Richard: Grundlagen der Bruchmechanik und des Ermüdungsrisswachstums. In: H. A. Richard, M. Sander: Ermüdungsrissausbreitung. Fachgruppe Angewandte Mechanik, Paderborn, 2003, Teil 1.
- [2] M. Sander: Experimentelle Ergebnisse von Überlast-, Blocklast- und Betriebslastexperimenten. In: H. A. Richard, M. Sander: Ermüdungsrissausbreitung. Fachgruppe Angewandte Mechanik, Paderborn, 2003, Teil 3.
- [3] 14. Entwurf zum UIC-Merkblatt 510-5 – Technische Zulassung von Vollrädern, Dezember 2001.
- [4] Entwurf DIN EN 13979-1: Bahnanwendungen, Radsätze und Drehgestelle, Räder, Technische Zulassungsverfahren – Teil 1: geschmiedete und gewalzte Räder. Beuth-Verlag, Berlin, 200.
- [5] M. Fulland, M. Schöllmann, H. A. Richard: Simulation of fatigue crack propagation processes in arbitrary three-dimensional structures with the program system ADAPCRACK3D. In: Ravi-Chandar et al. (Eds.): Advances in Fracture Research, CD-ROM Proceedings of the 10th International Conference on Fracture, Fatigue and Fracture, Honolulu, USA, 200.
- [6] H. A. Richard, M. Fulland, M. Schöllmann, M. Sander: Simulation of fatigue crack growth using ADAPCRACK3D. In: A. F. Blom (Ed.): Fatigue 2002, Proceedings of the Eighth International Fatigue Congress, 2002, Stockholm, Sweden, pp. 1405-1412.
- [7] M. Sander, H. A. Richard, G. Kullmer: ^{FAM}Control – Ein Mess- und Steuerungssystem zur automatischen Durchführung von Rissausbreitungsversuchen bei beliebiger Belastung. In: Frenz, H.; Wehrstedt, A. (Hrsg.): Kennwertermittlung für die Praxis, WILEY-VCH Verlag, Weinheim, 2002, S. 160-166.
- [8] M. Sander, H. A. Richard: Ermittlung von bruchmechanischen Kennwerten im Bereich der Verkehrstechnik. In: DVM-Bericht, Fortschritte der Bruch- und Schädigungsmechanik, Deutscher Verband für Materialforschung und -prüfung e. V., Berlin, 2004.
- [9] H. A. Richard, M. Sander, G. Kullmer, M. Fulland: Finite-Elemente-Simulation im Vergleich zur Realität. MP Materialprüfung 46 (2004) S. 441-448.

Warum lange suchen?

AOK Studenten-Service auf dem Campus.



Bei uns finden Sie ...

**Eine günstige
Krankenversicherung für
Studierende
mit tollem Service in zentraler
Lage.**

**AOK Studenten-Service Paderborn
Universität Paderborn
Gebäude ME 0 Raum 211**

**Telefon: 05251/124-424/-436
Fax: 05251/124-429
E-Mail: ASS.Paderborn@wl.aok.de**

**Mo, Mi & Do 10.00 Uhr bis 16.00 Uhr
Di 10.00 Uhr bis 17.30 Uhr
Fr 10.00 Uhr bis 13.00 Uhr
oder nach Vereinbarung**



www.aok.de



High-Tech-Unternehmensgründungen

Wer überlebt und wer überlebt nicht?

Forschungsnähe und eine enge Verbindung mit der Wissenschaft sind wichtige Erfolgsfaktoren für Unternehmensgründungen im High-Tech-Bereich. Diese Aussage lässt sich am Beispiel Paderborner Unternehmensgründungen empirisch erhärten. Sie bedarf aber einiger Ergänzungen und Verfeinerungen. Das Entstehen neuer Unternehmen, deren oftmals erfolgreiches Wirken und manchmal auch schnelles Ende mit dem Ausscheiden aus dem Wettbewerb, folgt Gesetzmäßigkeiten, die durch die Population-Ecology-Theorie (Hannan/Freeman 1989) nachvollzogen und interpretiert werden können.

Diese Variante der Organisationstheorie basiert auf evolutionstheoretischen Überlegungen, die in den Wirtschafts- und Sozialwissenschaften eine relativ lange Tradition haben. Zumindest lassen sich auf dieser Grundlage tragfähige Schlussfolgerungen ziehen und Ratschläge für adäquates unternehmerisches Verhalten in sich dynamisch entwickelnden Märkten wie denen des High-Tech- oder – spezifischer – des IT-Sektors ableiten.

Die Gründung der vielen Unternehmen im IT-Sektor in den zurückliegenden Jahren wurde von Beginn an von skeptischen Kommentaren begleitet. Der Tenor vieler Kommentare war: Das kann auf die Dauer nicht gut gehen. Die ersten Konkurse, denen bald weitere folgten, wurden als Bestätigung für diese skeptischen Prognosen angesehen. Was in den letzten Jahren im IT-Sektor geschah und in den nächsten Jahren geschehen wird, ist jedoch in einem marktwirtschaftlich organisierten Wirtschaftssystem völlig normal: Zunächst einmal gilt, dass laufend neue Unternehmen gegründet werden und ständig weniger erfolgreiche Unternehmen aus dem Markt ausscheiden. Im besonderen Fall des High-Tech-Sektors Informations- und Kommunikationstechnik verlief dieser Prozess im Übrigen geradezu lehrbuchmäßig, wobei der Population-Ecology-Ansatz der Organisationstheorie besonders gute Erklärungsleistungen liefert. Die Kernelemente dieser Theorie sollen deshalb im Folgenden skizziert werden.

Der Population-Ecology-Ansatz erklärt das Entstehen und Verschwinden von High-Tech-Unternehmensgründungen

Der Population-Ecology-Ansatz erklärt die Entstehung und Veränderung von Märkten und Organisationen als evolutionären Prozess. Analysegegenstand sind nicht die einzelnen Unternehmen, sondern Populationen von Unternehmen. Evolutionstheoretiker betonen die Bedeutung von Selektionsmechanismen für die Entwicklung von Organisationen. Während gut an die Umweltbedingungen angepasste Unternehmen überleben und kopiert werden, werden unterlegene Formen eliminiert.



Prof. Dr. rer. pol. habil., Dr. h.c. mult., Dipl.-Kfm. Wolfgang Weber ist seit 1985 Professor für Betriebswirtschaftslehre, insbes. Personalwirtschaft in der Fakultät für Wirtschaftswissenschaften. Er war von 1995 bis 2003 Rektor der Universität Paderborn. Seine Hauptarbeitsgebiete sind Unternehmensverhalten sowie Strategisches und Internationales Personalmanagement.

Die Entwicklung von Märkten, also von Organisationspopulationen, zeigt ein typisches Muster. Märkte wachsen in ihrer Entstehungsphase zunächst langsam, gehen dann in eine Phase rapiden Wachstums über, bevor die Anzahl der Unternehmen wieder sinkt und schließlich ein relativ konstantes Niveau erreicht (siehe Abbildung 1). Ursächlich für diesen Verlauf sind Wettbewerbs- und Legitimationsprozesse. Wettbewerb bezeichnet das Ausmaß, zu dem Unternehmen auf gleiche Ressourcen angewiesen sind. Legitimation beschreibt, inwieweit die Organisation sowie ihre Produkte und Dienstleistungen akzeptiert sind und als zuverlässig wahrgenommen werden.

Pioniere, die einen neuen Markt mit innovativen Produkten erschließen, haben kaum mit Wettbewerbsproblemen zu rechnen. Ihr Erfolg hängt nicht primär von der Effizienz ab, sondern von ihrer Fähigkeit, Akzeptanz für ihr Unternehmen und ihr Produkt bei Kapitalgebern und Kunden zu erlangen. Unternehmen, die sich kein hohes Maß an Akzeptanz sichern können, haben eine nur geringe Überlebenswahrscheinlichkeit.

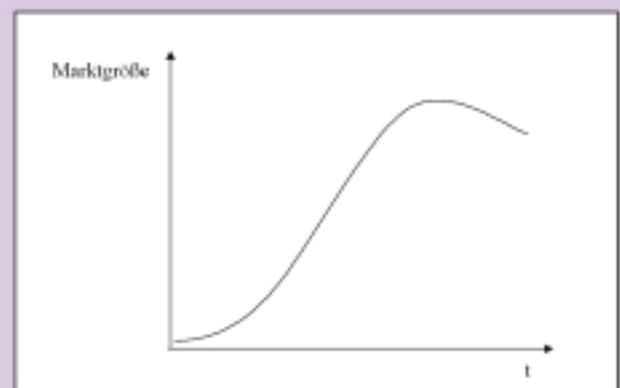


Abb. 1: Typischer Verlauf der Entwicklung eines neuen Marktes.

Viele Neugründungen scheiden in dieser Phase wieder aus dem Marktgeschehen aus, obwohl ihre Unternehmensidee innovativ und grundsätzlich tragfähig sein kann.

Den Pionieren folgen frühe Nachahmer, die die neuartigen Produkte imitieren oder mit Variationen auf dem noch jungen Markt auftreten. „Fast Mover“, die einen technologischen Vorsprung aufbauen oder wahren können, haben hier die besten Erfolgchancen. Die Effizienz der Leistungserstellung steht weiter hinter der Schnelligkeit zurück.

Mit zunehmender Anzahl erfolgreicher Unternehmen gewinnt ein Markt und gewinnen die in ihm agierenden Unternehmen an Legitimation. Viele möchten an dieser Möglichkeit erfolgreicher Unternehmenstätigkeit partizipieren, so dass die Unternehmenspopulation schnell und deutlich wächst. Die meisten Neugründungen orientieren sich an bestehenden, überlegenen Unternehmen, was dazu führt, dass erfolgreiche Unternehmen in besonders großem Maße imitiert werden und damit auch deren Unternehmenskonzept die größten Chancen hat, weiter verbreitet zu werden.

Je mehr Unternehmen auf einen Markt drängen, desto knapper werden verfügbare Ressourcen und desto höher wird die Wettbewerbsintensität. Die weniger effizienten Unternehmen werden eliminiert. Das bedeutet, dass die Population von Unternehmen in diesem stagnierenden Markt homogener wird, weil sich die erfolgreichen Unternehmenskonzepte bzw. Unternehmen schneller verbreiten und die weniger erfolgreichen Unternehmenskonzepte vom Markt verschwinden. Dabei tun sich jedoch Nischen auf, die ebenfalls gute Überlebenschancen für Neugründungen bieten.

Da in der Wachstumsphase mehr Unternehmen entstanden sind als der Markt dauerhaft erhalten kann, übersteigt die Zahl der Auflösungen die der Neugründungen, so dass der Markt insgesamt schrumpft.

Der Population-Ecology-Ansatz erlaubt es nicht nur, die Entwicklung des High-Tech-Sektors nachzuvollziehen, sondern ermöglicht auch die Identifikation von Variablen, die den Erfolg von Unternehmensgründungen beeinflussen. Der folgende Abschnitt dokumentiert dazu Ergebnisse einer empirischen Untersuchung von Unternehmensgründungen, die aus oder im direkten Umfeld der Universität Paderborn erfolgt sind (Schmelter 2004).

Hohes Wachstums-, Beschäftigungs- und Innovationspotenzial

Die meisten Impulse für den innovativen Strukturwandel gehen von Unternehmensgründungen in forschungs- und wissensintensiven Wirtschaftszweigen aus. Ihr Anteil an allen Gründungen in Deutschland ist nicht nur bis zum Ende der 90er Jahre auf ca. 25 Prozent gestiegen, ihnen schreibt man auch ein besonders hohes Wachstums- und Beschäftigungspotenzial zu. Hinweise darauf, inwiefern sich diese Erwartung empirisch bestätigen lässt und welche Faktoren den Erfolg forschungsorientierter Gründungen beeinflussen, liefert eine Befragung von (ehemaligen) Studierenden, wissenschaftlichen Angestellten und Professoren der Universität Paderborn, die ein Unternehmen gegründet haben und von Gründern, die sich im Technologiepark Paderborn und damit in unmittelbarer Nähe zur Universität angesiedelt haben. Die Stichprobe umfasst insgesamt 72 Unternehmen. Die Struktur der Stichprobe wird in Abbildung 2 dargestellt. 66

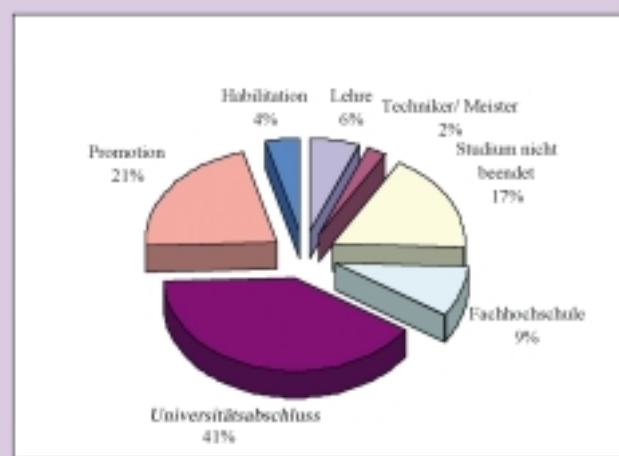


Abb. 2: Struktur der Stichprobe nach dem Abschluss der Gründer.

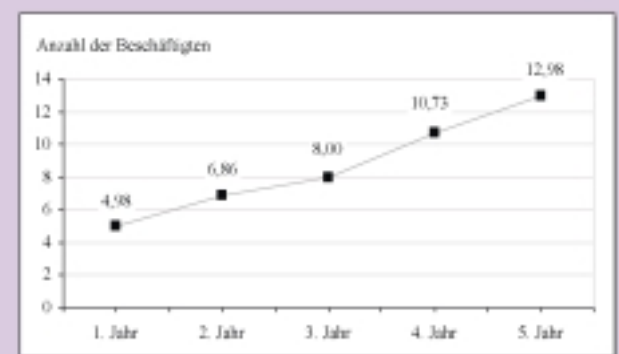


Abb. 3: Entwicklung der durchschnittlichen Beschäftigtenzahl.

Prozent der Gründer haben einen Universitätsabschluss, sind promoviert oder habilitiert. Eine zweite Gruppe wird durch die Unternehmen gebildet, deren Gründer einen Fachhochschulabschluss aufweisen, die das Studium nicht beendet haben oder eine andere Ausbildung angeben. Nicht enthalten sind in der Stichprobe Unternehmen, die keinerlei Forschungsbezug aufweisen, so dass ein Vergleich mit dieser Gruppe auf der Grundlage des verfügbaren Datenmaterials nicht möglich ist.

Erfasst wurde die Entwicklung dieser forschungsnahen Gründungen in den ersten fünf, besonders kritischen Jahren ihres Bestehens anhand der Mitarbeiterzahl und des Umsatzes.

Im ersten Jahr haben die Gründungen im Durchschnitt bereits fast fünf Beschäftigte (einschließlich der Gründer); bis zum fünften Jahr erhöht sich die Zahl auf fast 13 Mitarbeiter. Definiert man Betriebe mit einer durchschnittlichen Beschäftigungszunahme von mindestens zwei Arbeitskräften pro Jahr als „expansive Senkrechtstarter“, zählen 41 Prozent der untersuchten Unternehmen zu dieser Kategorie. Besonders wachstumsstark zeigen sich Gründungen, an denen Hochschulabsolventen, Promovierte oder Habilitierte beteiligt sind. Sie wachsen jährlich um 23 Prozent; die übrigen Gründungen jährlich nur um 6 Prozent. Ähnlich positiv gestaltet sich auch die Entwicklung der durchschnittlichen Jahresumsätze, wie der Abbildung 4 zu entnehmen ist.

Vom ersten bis zum fünften Geschäftsjahr erhöht sich der durchschnittliche Jahresumsatz auf das 2,4-fache. Differenziert man auch hier danach, welchen höchsten Ausbildungsabschluss die Gründer erworben haben, zeigt sich, dass die Unternehmen, die von Universitätsabsolventen, Promovierten bzw. Habilitierten

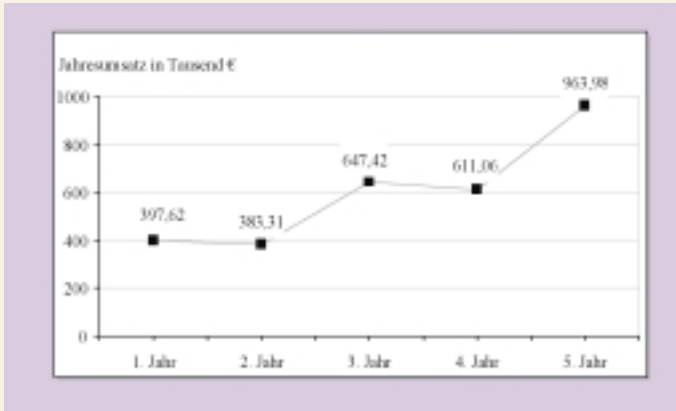


Abb. 4: Entwicklung des durchschnittlichen Jahresumsatzes in Tausend Euro.

gegründet wurden, ein höheres Umsatzvolumen aufweisen als die übrigen Unternehmen.

Diese Ergebnisse bestätigen das hohe Wachstumspotenzial innovativer und forschungsnaher Unternehmen und unterstreichen damit die besondere Bedeutung von Gründungen, die aus oder im direkten Umfeld von Hochschulen erfolgen.

Die meisten der hier untersuchten forschungsnahen Gründungen haben die Wachstumsphase des Marktes für ihren Unternehmensstart genutzt: Bei steigendem Marktwachstum wurden 39 Prozent, bei hohem Marktwachstum 26 Prozent und bei sehr hohem Wachstum 11 Prozent dieser 72 Unternehmen gegründet. Bei stagnierendem Wachstum stiegen 23 Prozent ein, bei sinkendem Wachstum begann nur ein Unternehmen seine Tätigkeit.

Bemerkenswert ist das Faktum, dass promovierte und habilitierte Gründer besonders häufig bei steigendem bis hohem Marktwachstum in eine Unternehmenstätigkeit einstiegen, während die Unternehmensgründungen von Technikern oder Meistern sowie Fachhochschulabsolventen bei niedrigerem Marktwachstum, also erst in der Phase eines starken Wettbewerbs erfolgten.

Der überwiegende Anteil der Unternehmensgründungen ordnet sich selbst als innovativ ein. Mit 65 Prozent ist der Anteil innovativer Gründungen relativ hoch. Eine Spezifizierung des Innovationsgrades ergibt, dass 34 Prozent ihr Produkt als neuartig einschätzen, während 19 Prozent ihr Produkt als Weiterentwicklung, zwölf Prozent als Abwandlung beziehungsweise Zusatzprodukt zu einem bestehenden Produkt deklarieren.

Die Untersuchung bestätigt die Annahme, dass es einen Zusammenhang zwischen dem Grad der Forschungsnahe und dem Innovationsgrad des ersten Produktes gibt. Universitätsabsolventen, Promovierte und Habilitierte erweisen sich nach eigener Einschätzung als die innovativeren Gründer. Die Hälfte von ihnen gibt an, ein neuartiges Produkt oder eine Weiterentwicklung eines bestehenden Produktes in den Markt eingeführt zu haben. Die Gründer, die maximal einen Fachhochschulabschluss haben, deklarieren ihr Produkt in 50 Prozent der Fälle als nicht innovativ. Insgesamt haben 56 Prozent der untersuchten Gründungen in den ersten fünf Jahren mit einer Forschungseinrichtung zusammengearbeitet, wobei Unternehmen, die mit der Universität Paderborn kooperieren (39 Prozent), oftmals auch mit anderen Universitäten, Fachhochschulen und Forschungsinstituten zusammenarbeiten.

Die Resultate unterstreichen die Bedeutung von Ausgründungen aus der Hochschule für den Transfer von wissenschaftlichem Know-how und Forschungsergebnissen in marktfähige Produkte und die Relevanz der Hochschule als wichtigen Impulsgeber für das Innovationspotenzial der regionalen Wirtschaft (Weber 2001).

Erfolgsfaktoren forschungsnaher Unternehmensgründungen

Um den Einfluss ausgewählter Faktoren auf die Entwicklung der Unternehmensgründungen in den ersten Jahren ihres Bestehens bestimmen zu können, wurden Regressionsanalysen – statistische Verfahren zur Überprüfung der Wirkung von unabhängigen Variablen auf eine Zielvariable – durchgeführt.

Forschungsnahe Gründungen wachsen der Untersuchung zufolge in den ersten Jahren dann am stärksten, wenn sie über ein hohes Startkapital verfügen, das Unternehmen in einer Phase hohen Marktwachstums entsteht, das Produkt nicht völlig neuartig ist und das Unternehmen mit einer Hochschule kooperiert.

Diese Ergebnisse bestätigen die populationsökologischen Überlegungen und ermöglichen die Ableitung von Handlungsempfehlungen für akademische Unternehmensgründer.

Als Pioniere setzen sich forschungsnaher Unternehmensgründer, die ein neuartiges Produkt vermarkten möchten, zwei spezifischen Problemen aus. Zum einen bedingt die Neuartigkeit die Zurückhaltung von Kunden und Kapitalgebern, in ein ihnen nicht bekanntes und möglicherweise nicht verständliches Produkt zu investieren. Neben diesen Problemen mangelnder Legitimation besteht die Gefahr, dass sie ihren Innovationsvorsprung gegenüber frühen Nachahmern verlieren, die aussichtsreiche Neuheiten identifizieren, verbessern und möglicherweise auf Kosten der Pioniere auf dem Markt durchsetzen. Diese Nachteile gegenüber den „Second Movers“, deren Wachstumspotential dieser Untersuchung zufolge höher ist, können die Pioniere durch die Kooperation mit einer Hochschule ausgleichen. Eine solche Zusammenarbeit kann die Legitimation der Neugründung erhöhen, indem das neue Unternehmen die Reputation der Forschungseinrichtung sowie deren Kontakte zu potenziellen Investoren und Kunden nutzt. Der Universität kommt im Rahmen des Kommunikationsprozesses zwischen den Beteiligten eine Multiplikatorfunktion zu, mittels derer eine Gründung ihre Akzeptanz erhöhen kann. Darüber hinaus erleichtert eine Forschungsk Kooperation die Sicherung des technologischen Vorsprungs gegenüber den Wettbewerbern. Sie wirkt sich offenbar auch positiv auf die schnellere Entwicklung forschungsintensiver Produkte zur Marktreife aus.

Eine weitere positive Konsequenz aus der Zusammenarbeit mit einer Universität ist der privilegierte Zugriff auf hochqualifizierte Fachkräfte, die insbesondere in sehr jungen innovativen Branchen häufig rar sind.

Bessere Entwicklungschancen als die Pioniere haben Unternehmensgründungen, die während der Expansionsphase der Marktentwicklung entstehen. Ihre Unternehmensform und ihr Produkt sind durch den Verweis auf (erfolgreiche) Vorgänger zu legitimieren, und sie können von den Fehlern der Pioniere lernen. Diese Phase der dynamischen Marktentwicklung ist durch Produkt- und Prozessvielfalt gekennzeichnet. Im Kampf um Standards, die sich letztlich durchsetzen, zählen insbesondere Schnelligkeit und Flexibilität. Deshalb zahlen sich Forschungsk Kooperationen in dieser Phase aus, wenn die Gründer durch sie Zugriff auf neue, überlegene Verfahren und Produkte erhalten, die ihnen im stärker werdenden Wettbewerb Vorteile verschaffen.

Mit zunehmender Wettbewerbsintensität gewinnt die Effizienz der Leistungserstellung an Bedeutung. Der Übergang von starkem Marktwachstum zur Phase der Stagnation des Marktes und damit zu steigender Wettbewerbsintensität stellt die zentrale Herausforderung für die jungen und innovativen Unternehmen



Abb. 5: Der TechnologiePark Paderborn liefert durch die Nähe zur Universität und die enge Verbindung der meisten Unternehmen mit der Universität ein hohes Maß an Akzeptanz auf dem Markt und Zugang zu hochqualifizierten Fachkräften.

dar: Je reifer ein Markt wird und je knapper damit Ressourcen werden, desto bedeutender werden effiziente Strukturen und Verfahren. Viele innovative Gründungen scheitern daran, dass ihnen die Umstellung von einem Innovationswettbewerb auf einen kostengetriebenen Wettbewerb nicht gelingt. Sie werden von Unternehmen, die lediglich Vorhandenes imitieren, aber mit geringeren Kosten arbeiten, vom Markt gedrängt.

Ein typisches Defizit der Gründungsforschung trifft auch diese Untersuchung. Es werden nur Unternehmen erfasst, die zum Befragungszeitpunkt noch bestehen. Für die im Technologiepark angesiedelten Gründungen lässt sich jedoch nachvollziehen, dass der überwiegende Teil der angeschriebenen Unternehmen noch immer existiert. Insbesondere ist keines der innovativen Unternehmen zwischenzeitlich ausgeschieden.

Grundlagenforschung schafft Innovationspotenzial

Aus den oben skizzierten Zusammenhängen und empirischen Belegen können Gründer und kommunale Wirtschaftsförderer, vor allem aber auch die Universität Schlussfolgerungen ziehen. Die Universität trägt durch die Unterstützung und Förderung von Unternehmensgründern, durch die Kooperation mit den Einrichtungen der Wirtschaftsförderung, durch einen engen Kontakt mit dem TechnologiePark Paderborn (Abbildung 5) zur wirtschaftlichen Entwicklung und Prosperität von Stadt und Region bei. Diesen Aufgaben kann deshalb gar nicht genug Aufmerksamkeit geschenkt werden. Sie liefern einen zentralen Beitrag zur Akzeptanz und Reputation der Universität als ein wichtiger Wirtschaftsfaktor. Je höher die Akzeptanz und Reputation der Universität auf diesem Gebiet ist, umso mehr verleiht sie Unternehmensgründern Legitimität, die zur Sicherstellung des Erfolgs insbesondere von frühen Innovatoren wichtig ist. Voraussetzung dafür, dass dieser wechselseitige Prozess der Produktion von Legitimität, Akzeptanz und Reputation sich im Zeitablauf verstärkt ist allerdings das „Kerngeschäft“ der Universität, die Produktion und Weitergabe von neuem Wissen durch Forschung und Lehre (Weber 2001).

Voraussetzung für eine forschungs- und technologiegetriebene wirtschaftliche Entwicklung ist die Existenz einer Kette von der Grundlagenforschung über angewandte Forschung und Entwicklung bis zur Produktentwicklung und der marktlichen Verwer-

tung dieser Produkte. Herausragende Bedeutung hat dabei das Fundament, auf dem diese Entwicklungsmöglichkeiten aufbauen, die Grundlagenforschung. Viele Beispiele zeigen, dass Innovationsprozesse sich verlangsamen oder sogar zum Erliegen kommen, wenn die Wissensanwender und Entwickler nicht von einer ausgebauten und erfolgreichen Grundlagenforschung zehren können. Sie schafft das innovative Forschungsumfeld, das die anwendungsnäheren Forschungsbereiche anregt.



Grundlagenforschung liefert die Basis für Innovationen.

Literatur

Aldrich, H. E. (1999): Organizations Evolving, London.

Fallgatter, M. (1999): Theorie des Entrepreneurship – Perspektiven zur Erforschung der Entstehung und Entwicklung junger Unternehmen, Wiesbaden.

Kieser, A./Woywode, M. (2002): Evolutionstheoretische Ansätze, In: Kieser, A. (Hrsg.): Organisationstheorien, 5. Aufl., Stuttgart/Berlin/Köln, S. 253-286.

Hannan, M. T./Freeman, J. (1989): Organizational Ecology, Cambridge MA.

Schmelter, A. (2004): Entwicklungsverläufe forschungsnaher Unternehmensgründungen und deren Determinanten, In: Die Betriebswirtschaft, Vol. 64, Nr. 4, S. 471-486.

Weber, W. (2001): Wissenschaft, neue Technologien und unternehmerisches Handeln, In: Blecker, Th. Gemünden, H. G. (Hrsg.): Innovatives Produktions- und Technologiemanagement, Berlin u. a., S. 589-601.



Dipl.-Kauffrau Anja Schmelter ist wissenschaftliche Mitarbeiterin in der Arbeitsgruppe Personalwirtschaft an der Universität Paderborn. Ihre Arbeitsgebiete sind soziale Netzwerke sowie Organisations- und Gründungsforschung.

Unsere Mitarbeiter –
eine Partnerschaft mit Perspektive

Hochschulabsolventen

- Ingenieure
- Wirtschaftswissenschaftler

Sie haben praxisorientiert studiert und stehen kurz vor dem Abschluss. Sie sind begeisterungsfähig und würden Ihre Ideen gern in einem internationalen Unternehmen kreativ umsetzen?

Wir bieten Ihnen die Herausforderung, die Sie suchen. In der Benteler-Gruppe werden Sie in innovativen Teams eingebunden und in unterschiedlichen Projekten gefordert. Ihre Konzepte zur Lösung von Problemen interessieren uns.

Sie sind noch im Studium? Kein Problem, auch ein Praktikum oder eine Diplomarbeit könnte der erste Einstieg sein.

Nutzen Sie Ihre Chance, wir freuen uns auf Ihre Bewerbung.

www.benteler.de



Take Your Chance

Machen Sie Ihre Karriere bei einem "Top 100 Unternehmen" der deutschen Industrie! Rund 18.000 Mitarbeiter beschäftigt die international aufgestellte Benteler-Gruppe. Mit den Bereichen Automobiltechnik, Stahl/Rohr, Maschinenbau und Handel zählen wir zu den Marktführern.

Benteler AG
Personalentwicklung
Sabine Peter
Residenzstraße 1
33104 Paderborn
Tel.: 0 52 54.81-18 46
konzern_pe@benteler.de

BENTELER 

Automobiltechnik • Stahl/Rohr • Maschinenbau • Handel

Viel Licht und keine Blendung

Aktive Scheinwerfersysteme zur Erhöhung der Verkehrssicherheit

Der Mensch nimmt ca. 90 Prozent der Informationen durch visuelle Wahrnehmung auf. Gute Sicht ist für das Führen eines Kraftfahrzeuges wichtig. Insbesondere bei Nacht und in der Dämmerung sind schlechte Sichtverhältnisse die Ursache von schweren Verkehrsunfällen. Statistische Untersuchungen haben gezeigt, dass sich 50 Prozent der tödlichen Verkehrsunfälle bei Nachtfahrten ereignen, obwohl in dieser Zeit nur 25 Prozent der durchschnittlichen Kilometerleistung zurückgelegt werden [Langwieder und Bäumler, 1997]. Dass dies hauptsächlich auf die schlechten Sichtverhältnisse und nicht auf andere Einflussfaktoren, wie z. B. Alkohol, Müdigkeit, o. Ä. zurückzuführen ist, wurde von [Sullivan und Flannagan, 2002] in einer umfangreichen Studie nachgewiesen.

Die Frage nach der optimalen Ausleuchtung des Verkehrsraumes beschäftigt die automobile Lichttechnik seit vielen Jahren. Einerseits sollen die Straße und ihr Umfeld möglichst hell ausgeleuchtet werden, damit der Fahrer die Objekte im Verkehrsraum sicher erkennen kann. Andererseits dürfen andere Verkehrsteilnehmer und der Fahrer selbst nicht geblendet werden. Dieser Zielkonflikt zwischen Ausleuchtung, Eigen- und Fremdblendung stellt, in der Terminologie der Ingenieurwissenschaften, den „zentrierenden Defekt“ dar, dessen möglichst gute Ausbalancierung oder gar Auflösung die wissenschaftliche Aufgabe des Ingenieurs ist. Die klassische und bis vor kurzem einzige Lösung besteht im Umschalten zwischen Fernlicht und Begegnungslicht, das umgangssprachlich oft als Abblendlicht bezeichnet wird. Das Fernlicht stellt die für die Ausleuchtung der Straße optimierte Lichtverteilung dar. Das Begegnungslicht ist sozusagen die Kompromisslösung, um Blendung zu vermeiden. Dass man indes auch von einem perfekt eingestellten Begegnungslicht geblendet werden kann, weiß jeder, dem schon einmal bei Nacht ein Fahrzeug an einer Bergkuppe entgegenkam.

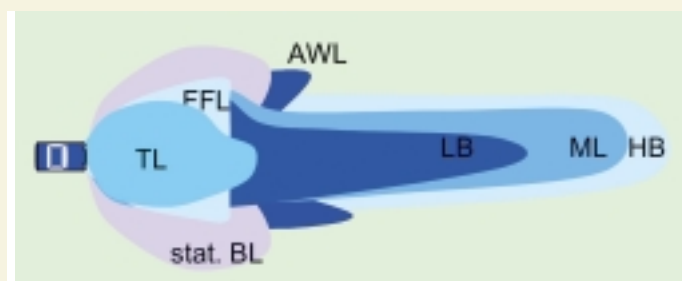
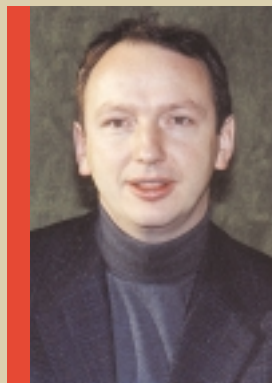


Abb 1: AFS-Lichtverteilungen: TL = Town Light (Stadtlit); FFL = Front Fog Light (Nebellicht); AWL = Adverse Weather Light (Schlechtwetterlicht); stat. BL = static Bending Light (statisches Kurvenlicht); LB = Low Beam (Abblendlicht); ML = Motorway Light (Autobahnlicht); HB = High Beam (Fernlicht).



Prof. Dr.-Ing. Jörg Wallaschek ist seit 1992 Professor für Mechatronik und Dynamik am Heinz Nixdorf Institut der Universität Paderborn. Seit 2001 ist er hochschulseitiger Vorstand des L-LAB, dem gemeinsamen Forschungsinstitut von HELLA KGaA Hueck und Co. und der Universität Paderborn. Seine Forschungsschwerpunkte sind piezoelektrische Systeme und Aktorik, Lichttechnik und mechatronische Systementwicklung.

Der Lösungsansatz: „Mechatronik hilft“

Auch in anderen Bereichen der Kraftfahrzeugtechnik gibt es ähnliche Zielkonflikte. So zum Beispiel bei der Abstimmung der Fahrwerksdämpfung. Für eine hohe Sicherheit müsste eine hohe Dämpfung, und für einen guten Komfort eine möglichst geringe Dämpfung gewählt werden. Auch dem ABS-System liegt ein Zielkonflikt zugrunde. Um eine möglichst starke Verzögerung des Fahrzeuges zu erreichen, sollen die Räder möglichst stark abgebremst werden. Andererseits dürfen die Räder nicht blockieren, weil sonst keine Seitenführung mehr gegeben ist und das Fahrzeug nicht mehr gelenkt werden kann. In den beiden hier genannten und in vielen weiteren Fällen, war es möglich, die Zielkonflikte mit Hilfe mechatronischer Systeme aufzulösen. Was liegt also näher, als die Methoden der Mechatronik auf den die Kraftfahrzeug-Lichttechnik dominierenden Zielkonflikt anzuwenden?

Einer der Autoren dieses Beitrages stellte vor 7 Jahren die provokative Frage „Was wäre, wenn ihr Scheinwerfer wüsste, wo Sie hinfahren?“ [Wallaschek, 1998]. Von vielen belächelt, machten sich damals eine kleine Forschergruppe am Heinz Nixdorf Institut auf den Weg, um diese Frage zu beantworten. Schnell wurde klar, dass in der Nutzung der damals gerade entstehenden Fahrzeug-Umfeld-Sensorik in Verbindung mit neuen optischen Konzepten ein enormes Verbesserungspotential für die Kraftfahrzeug-Lichttechnik liegt.

Im Folgenden wird zunächst der aktuelle Stand in der Entwicklung der Kraftfahrzeug-Lichttechnik beschrieben, der durch die Einführung adaptiver Scheinwerfersysteme gekennzeichnet ist. Diese Systeme werden auch abgekürzt als AFS (= Adaptive Frontlighting Systems) bezeichnet. Anschließend wird ein weit darüber hinaus gehendes aktives Lichtsystem vorgestellt, das im

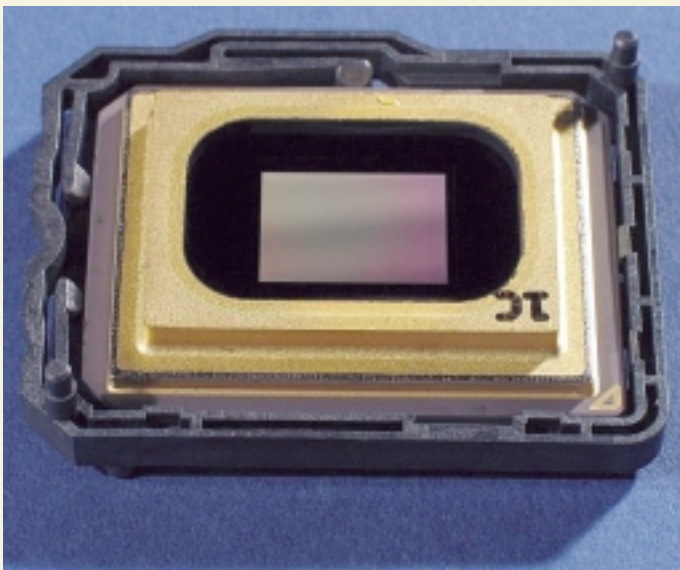


Abb. 2: DMD-Chip (Digital Micromirror Device-Chip), 10 x 14 mm² Gesamtspiegel-
gelfläche mit ca. 800 000 Mikrosiegeln.

L-LAB konzipiert und bereits als fahrfähiger Prototyp aufgebaut wurde.

Adaptive Scheinwerfersysteme

Für die Sicherheit bei Nachtfahrten stellen das heutige Fernlicht und Begegnungslicht nicht die nach dem Stand der Technik optimal mögliche Lösung dar [Huhn, 1999]. Eine einfache, nahe-
liegende Verbesserung besteht darin, das Fahrzeug mit speziellen
zusätzlichen Scheinwerfern auszurüsten, die vom Fahrer, je nach
Situation zu- oder wieder abgeschaltet werden. Ein nächster
Schritt besteht darin, diese zusätzlichen Lichtfunktionen nicht in
einzelnen Zusatzscheinwerfern zu realisieren, sondern sie in den
Hauptscheinwerfer zu integrieren und die Umschaltung
zwischen den jeweiligen Lichtverteilungen automatisch vorzu-
nehmen. Dies ist das Grundkonzept der adaptiven Scheinwerfer-
systeme.

Adaptive Scheinwerfersysteme ermöglichen es, diskret vordefi-
nierte Lichtverteilungen zu erzeugen. Abbildung 1 zeigt die
dabei vorgesehenen Lichtverteilungen, deren gesetzliche Zulas-
sung bis 2006/2007 erfolgen soll. Die Anpassung der Lichtver-
teilung erfolgt in Abhängigkeit von der Fahrzeuggeschwindigkeit,
dem Straßentyp und den Witterungsbedingungen, was eine
Verbesserung der heutigen Kraftfahrzeug-Lichttechnik darstellt.
Der zugrunde liegende „zentrierende Defekt“, nämlich der
Konflikt zwischen Ausleuchtung und Blendung, kann damit
aber nicht aufgelöst werden.

Aktive Scheinwerfersysteme – Die Idee

„Was wäre wenn ihr Scheinwerfer nicht nur wüsste, wo Sie
hinfahren, sondern auch welche anderen Verkehrsteilnehmer sich
im Verkehrsraum befinden?“

Die Grundidee der aktiven Scheinwerfersysteme besteht in der
ersten Stufe darin, ständig mit Fernlicht zu fahren, dieses aber so
zu manipulieren, dass es nur die Bereiche des Verkehrsraumes
erreicht, in denen es keine Blendung verursacht. In Deutschland
erfolgen bisher nur 5 Prozent aller Nachtfahrten mit Fernlicht
[Hella Research Review, 2000]. Die Ersetzung des suboptimalen

Begegnungslichtes durch das Fernlicht, das sehr gute Sichtbedin-
gungen gewährleistet, würde dann bereits zu einer wesentlichen
Steigerung der Verkehrssicherheit und des Fahrkomforts führen.

Aktive Scheinwerfersysteme – Die Konzipierung

Am Beginn der Entwicklung eines mechatronischen Systems
steht die Konzipierung der logischen Systemarchitektur, die auch
funktionaler Entwurf genannt wird. Die vom System zu erfüllen-
den Aufgaben und die dabei gestellten Anforderungen werden
systematisch erfasst und analysiert. Am Ende dieses Prozesses
steht die Funktionsstruktur, die eine modulare hierarchische
Beschreibung der vom System zu erfüllenden Aufgaben darstellt.
Funktionsstrukturen werden lösungsneutral formuliert, d. h., es
spielt noch keine Rolle, mit welchen physikalischen Effekten
und maschinenbaulichen Lösungen die Aufgabe später tatsäch-
lich gelöst wird.

Das übergeordnete Ziel unseres Projektes ist die Schaffung eines
Scheinwerfersystems, das eine möglichst hohe Erkennbarkeitsent-
fernung hat und dabei eine vorgegebene maximale Beleuchtungs-
stärke in den Augenpunkten der anderen Verkehrsteilnehmer
nicht überschreitet.

Das „blendungsfreie Fernlicht“ wurde im Rahmen eines

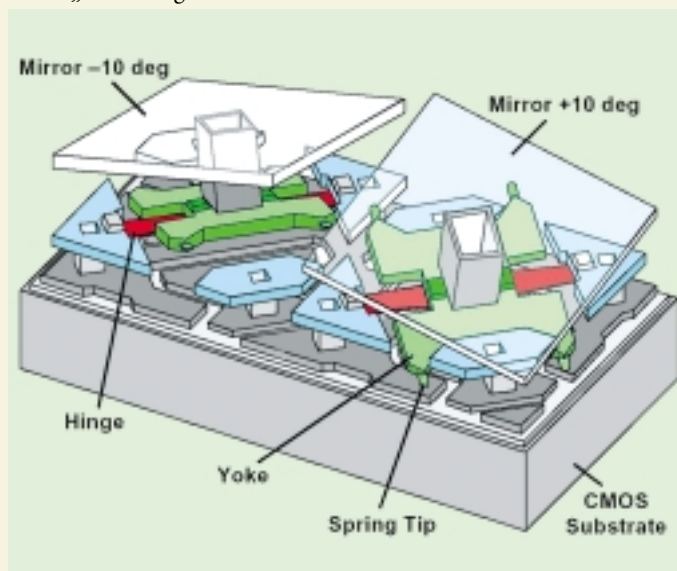


Abb. 2a: DMD-Chip, 2 Mikrospiegel, die um ein Torsionsgelenk (rot) elektro-
statisch in eine der beiden Endlagen bewegt sind. [Texas Instruments].



Abb. 2b: DMD-Aktiver Scheinwerfer integriert in ein Serien-Scheinwerfer-Gehäuse.

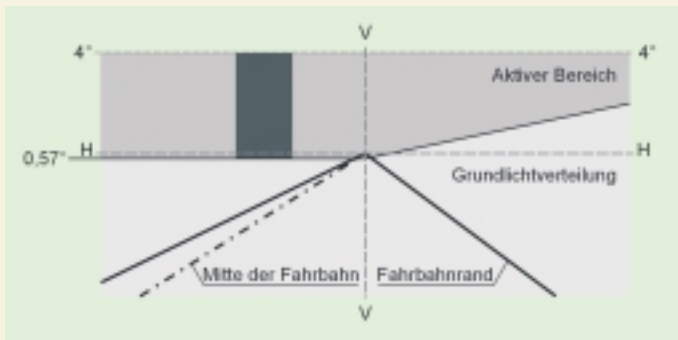


Abb. 3: Prinzipdarstellung der einzelnen Bereiche der Lichtverteilung in einem Messschirm nach ECE.

Forschungsprojektes der „International Graduate School of Dynamic Intelligent Systems“ von [Roslak, 2003] untersucht. Die Lichtverteilung im Verkehrsraum wird als Überlagerung der individuellen Lichtverteilungen aller am Verkehrsgeschehen beteiligten Objekte betrachtet und soll so beeinflusst werden, dass die Wahrnehmungsbedingungen für alle am Geschehen beteiligten Verkehrsteilnehmer optimal sind. Dabei werden Satelliten-Navigation, fahrzeugeigene Sensorik und Fahrzeug-Fahrzeug-Kommunikation genutzt, um Position, Orientierung und Geschwindigkeit der Verkehrsteilnehmer zu erfassen.

Bei der Schaffung optimaler Sichtbedingungen ist es notwendig, die Straßentopologie zu berücksichtigen. Über die Funktion des „dynamischen Kurvenlichtes“ hinaus, hat [Kuhl, 2004] untersucht, wie die Lichtverteilung auch an Kuppen und Wannen im Straßenverlauf angepasst werden kann.

Eine weitere Systemfunktion, die im Rahmen von Studien- und Diplomarbeiten untersucht wurde, ist das „fahrerspezifische Licht“, mit dem der Fahrer im Rahmen der gesetzlichen Vorgaben die Möglichkeit erhält, die Lichtverteilung seines Fahrzeuges in physiologisch und verkehrssicherheitstechnisch sinnvollen Ausprägungen individuell anzupassen. Variationsparameter können dabei die Schärfe der Hell-Dunkel-Grenze, der Leuchtdichtegradient des Vorfeldes, die Streubreite der Scheinwerfer oder die Helligkeit der eigenen Fahrbahn sein.

Über die bereits genannten Systemfunktionen hinaus wurde im Rahmen der Systemkonzipierung auch das „Markierungslicht“, zur Hervorhebung gefährdeter Objekte und Verkehrsteilnehmer in Gefahrensituationen, und das „Displaylicht“ untersucht, bei dem die Fahrbahnoberfläche als Projektionsfläche verwendet wird, auf der z. B. Navigationshinweise oder Geschwindigkeiten eingeblendet werden.

Aktive Scheinwerfersysteme – Die Umsetzung

Es ist klar, dass die zur Umsetzung dieser Konzeptionen erforderlichen Technologien sehr vielfältig sind. Schon eine Darstellung der für die Informationsgewinnung denkbaren Lösungsmöglichkeiten würde den Rahmen sprengen. Deshalb werden hier nur zwei Konzepte beschrieben. Weiterführende Darstellungen können [Roslak, 2004] entnommen werden. Eine erste Lösung basiert ausschließlich auf der Nutzung fahrzeugbasierter Sensorik. Dabei werden Sensorsysteme wie z. B. Radar, Lidar, Infrarot-Scanner o. Ä. genutzt um Informationen über das Fahrzeugvorfeld und die -umgebung zu gewinnen. Diese Lösung ermöglicht eine lokal autonome Betrachtung des Verkehrsgeschehens vom individuellen Standpunkt der daran beteiligten Fahrzeuge. Ein komplett anderes Lösungsprinzip besteht in der

Verwendung einer Inter-Fahrzeug-Kommunikation. Hier wird jedes Fahrzeug als Teilnehmer eines Netzwerkes betrachtet. Durch Informationsaustausch der Fahrzeuge untereinander werden die zur Ansteuerung des Lichtsystems erforderlichen Informationen generiert. Das Kommunikationsnetz stellt damit die Basis für die Erkennung von Verkehrssituationen dar.

Um die beim aktiven Licht geforderte völlig flexible Lichtverteilung erzeugen zu können, musste ein völlig neues lichttechnisches System aufgebaut werden. Dies geschah im Rahmen einer Dissertation im L-LAB [Kauschke, 2004]. Die Wahl zwischen verschiedenen Optikkonzepten fiel auf den so genannten DMD-Scheinwerfer. Abbildung 2 zeigt eine Prinzipdarstellung dieses Systems. Der Kern besteht aus einem DMD-Chip (DMD = Digital Micromirror Device), einem Mikrospiegelarray, dessen Einzelspiegel pixelweise adressiert werden können. Dadurch kann die Lichtstärke einzelner Bildpunkte durch Pulsweitenmodulation gezielt eingestellt werden. Wesentliche Herausforderungen bei der Entwicklung des DMD-Scheinwerfers bestanden in der Einkopplung des von der Kurzbogenlampe ausgesandten Lichtstroms auf den DMD-Chip und die anschließende Projektion des DMD-Chips durch die Auskoppeloptik auf die Straße. Es mussten erhebliche Leistungssteigerungen gegenüber konventionellen Video-Projektoren erzielt werden, um die geforderten Werte zu erreichen.

Die Lichtverteilungen werden anhand der in Abbildung 3 gezeigten Ausleuchtung eines Messschirmes nach ECE charakterisiert. Das unmittelbare Fahrzeugvorfeld wird durch eine Basislichtverteilung, die dem heutigen Abblendlicht entspricht, permanent ausgeleuchtet. Auf diese Weise wird das komplette Abschatten eines auch nicht selbst leuchtenden Objektes verhindert. Die Lichtverteilung im Bereich oberhalb der Hell-Dunkel-Grenze wird aktiv verändert, indem z. B., wie in Abbildung 3 gezeigt, ein Streifen (dunkelgrau gekennzeichnet) ausgeblendet wird, wenn sich dort der Augenpunkt eines anderen Verkehrsteilnehmers befindet. Aufgrund von Analysen der Verkehrsraumgeometrie und der Fixationsstellen wurde die obere Grenze des aktiven Feldes auf 4° oberhalb der Basislichtverteilung begrenzt.

Das von uns entwickelte aktive Scheinwerfersystem wurde in einem Versuchsfahrzeug implementiert. Als Testplattform wurde ein Fahrzeug vom Typ Volvo V70 mit einem angebauten Front-Rack verwendet. Die wesentlichen Komponenten des in Abbildung 4 gezeigten Systems sind:

- Vorausschauende Sensorik (LIDAR) zur Erfassung der Position und Differenzgeschwindigkeit der Fahrzeuge im Vorfeld.
- Fahrzeugsensorik zur Ermittlung der eigenen Geschwindigkeit und des aktuellen Radius der Fahrbahn.
- DMD-Scheinwerfer zur Erzeugung von variablen Lichtverteilungen (DMD – Digital Micromirror Device).
- CAN-Bus zum Informationsaustausch.
- Steuerrechner.

Experimentelle Ergebnisse

Das in dem Versuchsfahrzeug implementierte aktive Scheinwerfersystem wurde bereits unter realen Verkehrsbedingungen erprobt. Das Ziel der Untersuchung war dabei, die Eigenschaften dieses lichttechnischen Systems unter technischen und psychologischen Gesichtspunkten zu bewerten.

Eine wesentliche Beurteilungsaufgabe betraf die Einschätzung



Versuchsfahrzeug mit aktivem Scheinwerfersystem.

des potentiellen Nutzens des Systems im nächtlichen Straßenverkehr. Im Mittelpunkt standen hier Fragestellungen bezüglich der Verbesserung der Erkennbarkeitsbedingungen und möglicher Einflüsse der dynamischen Veränderung der Lichtverteilung auf den Wahrnehmungsprozess des Systemnutzers.

Die Einschätzung des Einflusses der aktiven Lichtverteilung auf Erkennbarkeitsverhältnisse anderer Verkehrsteilnehmer bildete die zweite Bewertungsaufgabe. Ein wesentlicher Punkt war hier die Bestimmung der Blendung anderer Fahrer. Die Untersuchung sollte zeigen, dass diese durch das aktive Lichtsystem vollständig ausgeschlossen werden kann.

Die Ermittlung der Beleuchtungsstärke im Augenpunkt des entgegenkommenden Fahrzeuges erfolgte mit Hilfe eines digitalen Beleuchtungsstärkemessgerätes dessen Photometerkopf auf der Höhe der Fahreraugen positioniert und mittig zur eigenen Fahrbahnspur gerichtet wurde. Die Messungen wurden für drei unterschiedliche Lichtverteilungen des aktiven Scheinwerfersystems durchgeführt. Neben dem aktiven Modus, mit eingeschaltetem System, wurde zum Vergleich auch die Beleuchtungsstärke für das aus dem gleichen Scheinwerfer erzeugte Abblend- und Fernlicht gemessen, siehe Abbildung 5. Die Ergebnisse zeigen, dass die vom aktiven Licht erzeugte Beleuchtungsstärke im Messpunkt den gesetzlich zugelassenen Grenzwert von 0,5 lx nicht überschreitet, solange der Abstand der Fahrzeuge im Messbereich der Sensoren ist, der hier zwischen 180 m und 30 m liegt.

Bei der Beurteilung der Ergebnisse ist zu beachten, dass das aktive Scheinwerfersystem mehr bietet als eine Umschaltung zwischen Fern- und Abblendlicht, denn außerhalb des abgedunkelten Bereichs der Augenpunkte des Gegenverkehrs wird ständig die Ausleuchtung des Fernlichts gegeben, siehe Abbildung 6.

Eine wesentliche Bewertung ist die subjektive Beurteilung des aktiven Lichtsystems durch nicht an der Entwicklung beteiligte Testpersonen. Insgesamt 21 Probanden hatten Gelegenheit, ein entsprechend ausgerüstetes Versuchsfahrzeug zu fahren. Nachtfahrten wurden mit Abblendlicht, Fernlicht und Aktivem Licht unter gleichen Verkehrsbedingungen durchgeführt. Der visuelle Eindruck wurde einmal aus der Perspektive des Fahrers des Fahrzeuges mit dem aktiven Lichtsystem und einmal aus der Perspektive des Fahrers des entgegenkommenden Fahrzeuges beurteilt. Nach jeder Fahrt wurde ein Fragebogen von den Versuchspersonen ausgefüllt.

Die vom aktiven Licht erzeugte Blendung der entgegenkommenden Fahrer wurde als kaum wahrnehmbar eingeschätzt. Dies gilt



Abb. 4: Komponenten des Versuchsfahrzeugs.

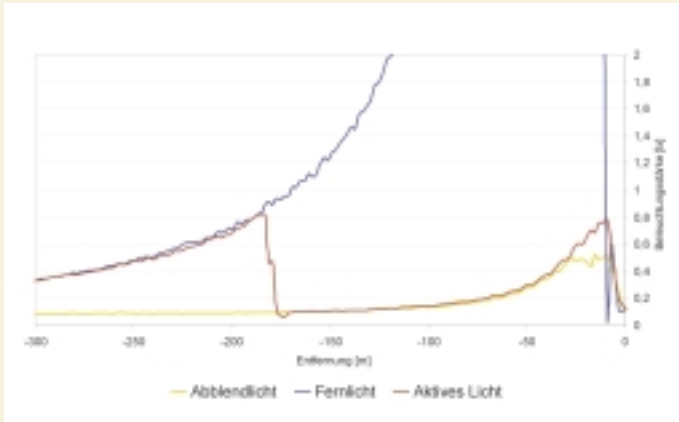


Abb. 5: Beleuchtungsstärke im Augenpunkt eines entgegenkommenden Fahrers für das Abblend-, Fern- und Aktive Licht.

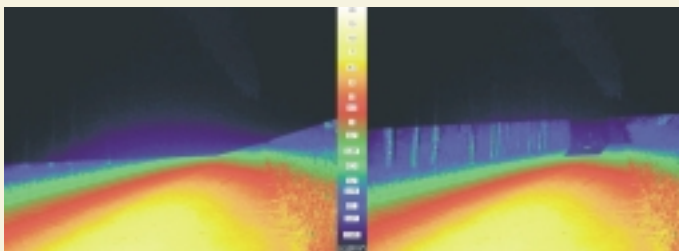


Abb. 6: Leuchtdichtebild des Fahrzeugvorfelds bei Abblendlicht (links) im Vergleich zu Aktivem Licht (rechts).

ebenfalls für Fahrzeuge im Kolonnenverkehr, bei denen reflektiertes Licht nachfolgender Fahrzeuge im Außenspiegel störend wirken kann. Die Anhebung der Hell-Dunkel-Grenze um 4° oberhalb des Abblendlichts führte zu einer erheblichen Steigerung der Erkennbarkeitsentfernung und dadurch der Sicherheit. Das System arbeitet mit einer hohen Zuverlässigkeit, auch in sehr komplexen Verkehrssituationen, wie z. B. bei einem großen Verkehrsaufkommen und bei hohen Differenzgeschwindigkeiten zwischen dem eigenen und einem anderen Fahrzeug.

Ausblick

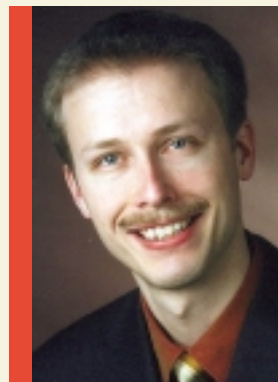
Aktive Scheinwerfersysteme stellen einen wesentlichen Schritt zur Weiterentwicklung der Kraftfahrzeug-Lichttechnik dar. Mit ihnen kann der Konflikt zwischen hoher Beleuchtungsstärke und geringer Blendung aufgehoben werden. Sie führen zu einer Steigerung der Verkehrssicherheit und erhöhen den Komfort.

Doch für die Lichttechniker und Mechatroniker im L-LAB ist die Entwicklung damit nicht am Ende. Neben der Untersuchung der Einsatzmöglichkeiten der vielen neuen optischen Technologien, wie z. B. anorganische und organische Leuchtdioden, die in den nächsten Jahren verfügbar werden, beschäftigt uns dabei vor allem eine Frage: „Was wäre, wenn ihr Scheinwerfer nicht nur wüsste, wo Sie hinfahren und welche anderen Verkehrsteilnehmer sich im Verkehrsraum befinden, sondern wenn er auch wüsste, wo Sie hinschauen und wie Sie sich fühlen?“

Literatur

- [Hella Research Review, 2000] Hella Lichttechnik Research & Development Review 2000.
 [Huhn, 1999] Huhn, W.: Anforderungen an eine adaptive Lichtverteilung für Kraftfahrzeugscheinwerfer im Rahmen der ECE-Regelungen; Dissertation, TH Darmstadt, 1999.
 [Kauschke, 2004] Kauschke, R.; Eichhorn, K.; Wallaschek, J.:

- Adaptive Scheinwerfer – Technologieüberblick, Licht 2004-Tagung, Dortmund, 19.-22.09.2004, Tagungs-CD-ROM.
 [Kuhl, 2004] Kuhl, P.: Heutige und zukünftige Kfz-Lichtfunktionen basierend auf vollbeweglichen Scheinwerfermodulen. In Proceedings zur Tagung „Intelligente Mechatronische Systeme“ 2004, HNI-Verlagsschriftenreihe, Nr.145, S.197-206.
 [Langwieder u. Bäumler, 1997] Langwieder, K.; Bäumler, H.: Charakteristik von Nachtunfällen. In Proceedings zur Tagung „Progress in Automobile Lighting“, TH Darmstadt 1997.
 [Roslak, 2003] Roslak J. ; Wallaschek J.: Aktive Kfz-Lichtverteilungen zur kollektiven Ausleuchtung des Verkehrsraumes. In Proceedings zur Tagung „Intelligente Mechatronische Systeme“ 2003, HNI-Verlagsschriftenreihe, Nr.122, S.29-38.
 [Roslak, 2004] Roslak, J.: Entwicklung eines Systems zur kollektiven Ausleuchtung des Verkehrsraums. Dissertation, Universität Paderborn, 2004.
 [Sullivan und Flannagan, 2002] Sullivan, J.M.; Flannagan, M.J.: The role of ambient light level in fatal crashes: inferences from daylight saving time transitions. Accident Analysis and Prevention 34 (2002), S. 487-498.
 [Wallaschek, 1998] Wallaschek, J.: Innovation – Weiß Ihr Scheinwerfer, wo Sie hinfahren? Hella Lichttechnik Research & Development Review 1998.



Dipl.-Ing. Rainer Kauschke studierte Maschinenbau, Vertiefung Lichttechnik, Konstruktion und Mechatronik seit 1996 an der TU Ilmenau. Ein Auslandsstudienjahr an der Universität Nottingham, England, ergänzte sein Studium. Von 2001 bis 2004 war er wissenschaftlicher Mitarbeiter an der Universität Paderborn, L-LAB, Fachgebiet Mechatronik und Dynamik, und forschte an aktiven Scheinwerfersystemen. Seit 10/2004 ist er Mitarbeiter bei der Hella KGaA Hueck & Co. in Lippstadt.



Dipl.-Ing. Jacek Roslak studierte Maschinenbau mit Vertiefungsrichtung Fahrzeugtechnik an der Technischen Universität Stettin (Polen). Während des Studiums absolvierte er ein einjähriges Stipendium der Carl-Duisberg-Gesellschaft in NRW. Seit 2001 ist er Stipendiat der Graduate School of Dynamic Intelligent Systems und wissenschaftlicher Mitarbeiter in der Fachgruppe Mechatronik und Dynamik an der Universität Paderborn. Im Rahmen seiner Dissertation bearbeitet er das Projekt Aktives Licht.



Werben in Paderborn

Wir gestalten Ihren Auftritt

Die PADA Werbeagentur
bietet Ihnen rund um Ihr Produkt
oder Dienstleistung,
eine auf den Markt gerichtete
offensive Werbung.

Dabei realisieren wir für Sie
vom Internetauftritt über
Logoentwicklung, Geschäftspapiere,
Zeitschriften, Prospekte, Flyer,
Grafiken, Poster bis zum
3D-Modelling alles was Sie
zur Werbung benötigen.

P A D A

Kreatives Handels - Marketing

Martin Heynen • Heierswall 2 • 33098 Paderborn
Tel.: 0 52 51/52 75 77 • FAX: 0 52 51/52 75 78
E-mail: pada-werbeagentur@t-online.de • www.pada-werbeagentur.de



Mobile-Banking: Stecken Sie Ihre Bank in die Tasche!

Die ideale Bankverbindung? Eine, über die Sie zu jeder Zeit an jedem Ort bequem Ihre Geldgeschäfte regeln.

Lassen Sie Ihr Online-Banking doch einfach von der Leine – holen Sie sich unsere PDA-Edition für den sicheren Kontakt zu Konto und Depot!

***Wir machen
den Weg frei***



**Volksbank
Paderborn-Höxter**
mit uns zum Erfolg