

III. Prozesse

Eine zentrale Aufgabe der Innovationsforschung liegt darin, die Verlaufsdynamik zu beschreiben und zu erklären, in der Neuerungen hervorgebracht und aufgenommen werden. Im Vordergrund steht hierbei die Identifizierung von Entwicklungsmustern, jener Strukturveränderungen, die wir Innovation nennen. Den abstrakten Hintergrund solcher Prozessbeschreibungen geben Modelle ab, welche die empirisch beobachteten Entwicklungsmuster in einen Erklärungsrahmen einbinden. Um bei all den Modellen und Theorien nicht den Überblick zu verlieren, ist es sinnvoll, darauf zu achten, auf welche Ebene der Beobachtung sich die theoretischen Aussagen beziehen. Zwei Ebenen können hierbei auseinander gehalten werden: die Ebene des gesellschaftlich-technologischen Wandels und die Ebene des organisierten Innovationsprozesses. Unternehmen gelten dabei als strategische Orte der Produktion von Innovation. Auf der Ebene des *gesellschaftlich-technologischen Wandels* werden Aussagen darüber getroffen, welche Entwicklungen außerhalb des Kontrollsektors eines einzelnen Unternehmens oder einer einzelnen Organisation ablaufen. Ein Unternehmen kann und muss sein Entscheidungsverhalten an diesen Entwicklungslinien ausrichten; bestimmen kann es sie nicht. Im Kontrast dazu bildet die Ebene des *organisierten Innovationsprozesses* den Referenzrahmen für Aussagen darüber, wie Innovationen innerhalb der Grenzen einer Organisation verlaufen und – da es immer auch um die Optimierung von Entscheidungsverhalten geht – verlaufen *sollen*. Während die erstgenannte Ebene Erkenntnisse liefert, an denen einzelne Akteure ihr politisches und wirtschaftliches Handeln orientieren können, wird auf der letztgenannten Managementwissen produziert, das als Handlungsmodell zur Organisation von betrieblichen Innovationsprozessen verwendet werden kann.

Neben dieser grundsätzlichen Ebenendifferenzierung können je nach Erklärungslogik zwei Modelltypen unterschieden werden. *Lineare* Modelle beschreiben den Ablauf als eine gerichtete Ursache-Folge-Sequenz, *non-lineare* dagegen heben zwei weitere Mechanismen hervor: Zum einen Rekursionen, wenn die zeitliche und sachliche Abfolge von Ursache und Wirkung umgedreht

wird, zum anderen Unterbrechungen im Laufe eines Innovationsprozesses, wenn es zu plötzlichen Strukturveränderungen und Rückschlägen kommt, die mit der Idee des gerichteten Innovationsverlaufs nicht vereinbar sind. Wie in Tabelle 1 zu sehen ist, möchte ich zunächst die linearen Modelle vorstellen, um dann separat auf die non-linearen einzugehen.

Tabelle 1: Ebenen und Modelle der Innovationsforschung

	Lineare Modelle	Non-lineare Modelle
Ebene des gesellschaftlich-technologischen Wandels	Technologieschub-Modell <i>(technology push)</i> Nachfragesog-Modell <i>(demand pull)</i>	Wellen-Zyklen-Modell Evolutionsmodell
Ebene des organisierten Innovationsprozesses	Phasen-Modelle	Ketten-, Reise-, Feuerwerk-Modelle

1. LINEARE MODELLE

Auch wenn sie später kritisiert werden, gehören die linearen Modelle zum Grundwissen der Innovationsforschung. Fast alle späteren Modelle bauen darauf auf, sodass es unerlässlich ist, deren Argumentation wiederzugeben.

1.1 Technologieschub und Nachfragesog als Ursache für Innovationen

Bei der Frage, welche fundamentalen Kräfte den gesellschaftlich-technologischen Wandel bestimmen, konkurrieren zwei Erklärungsmodelle miteinander. Das eine Modell hält die gesellschaftliche Nachfrage (*demand pull*), das andere das Angebot von neuer Technologie (*technology push*) für den treibenden Faktor des ge-

sellschaftlich-technologischen Wandels. Beiden Modellen liegt implizit ein Gleichgewichtsmodell zugrunde, in dem der gesellschaftlich-technologische Wandel dafür sorgt, ein Ungleichgewicht zwischen Angebot und Nachfrage zu beseitigen. Obgleich sie in Reinform in der Innovationsforschung kaum mehr vertreten werden, bilden diese beiden theoretischen Idealtypen nichtsdestotrotz die Ausgangsbasis für viele neuere Prozessmodelle. Außerdem lässt sich erst mit der Kenntnis dieser Modelle abschätzen, worin die Besonderheiten und die Leistungen der non-linearen Ansätze liegen.

Das *Nachfragesogs-Modell* geht davon aus, dass ein Anstieg des Haushaltseinkommens eine Nachfrage nach neuen und verbesserten Produkten entstehen lässt. Durch Marktforschung und durch den Preisauftrieb für innovative Produkte wird den Unternehmen signalisiert, aus diesem durch die Nachfrage induzierten Ungleichgewicht Nutzen zu ziehen, indem sie solche Innovationen entwickeln und anbieten, die diese Nachfrage stillen. Die Einsicht, dass die Nachfrage Innovationen stimuliert, wurde bahnbrechend von Jacob Schmookler (1966) formuliert und mit historischem Material belegt. So illustriert er seine Theorie zum Beispiel mit der Geschichte der Photographie, die Schmookler eng mit dem Aufstieg eines neuen wohlhabenden Bürgertums in Europa und vor allen Dingen in den USA verknüpft sieht. Der Habitus des Adels, sich auf Ölgemälden porträtieren zu lassen, wurde von der neuen bürgerlichen Elite kopiert. Doch der Mangel an Zeit, an zu verschwendemdem Kapital und nicht zuletzt das aufgrund der gestiegenen Nachfrage nach Porträtgemälden knappe Kontingent an ausgewiesenen Malern wurde zum Auslöser für die Innovation und Verbreitung der Studiophotographie. Nicht der Erfindung der Photographie, die ja bereits in Form der »camera obscura« seit Jahrhunderten in Künstler- und Wissenschaftlerzirkeln bekannt war, sondern der Nachfrage, die sich in Folge einer soziodemographischen Umschichtung der Gesellschaft abzeichnet, sei die Innovation der Studiophotographie zuzurechnen (ebd.: 93). Dem Modell des Nachfragesogs liegt ein kausal-chronologisches Muster zugrunde. Am Anfang steht die gesellschaftliche Nachfrage, die durch Marktsignale vermittelt Produzenten anzieht (*pull*), Innovationen zu entwickeln. Der

Markt fungiert hierbei als Entscheidungsinstanz darüber, ob ein neues Produkt zum Erfolg wird oder scheitert. Daher ist es für den Produzenten notwendig, im Vorhinein Bescheid zu wissen, in welche Richtung sich die Bedürfnisse der Konsumenten entwickeln. Die Möglichkeit, prognostisch Kenntnisse über die Konsumentennachfrage zu gewinnen, ist dann gegeben, wenn eine Innovation bereits erfolgreich adaptiert wurde, so dass aufgrund von Nutzungserfahrungen im konkreten Anwendungskontext offenkundig wird, wie dieses Produkt weiter verbessert und durch Modifikationen an Kundenwünsche angepasst werden kann.

Das Nachfragesog-Modell vermag Verbesserungsinnovationen zu erklären. Bei den ersten Notebooks war zum Beispiel das Gerät nur unzureichend lange im Akkubetrieb nutzbar; es entstand eine Nachfrage, die eine Reihe von Verbesserungsinnovationen stimulierte: Chip-Hersteller investierten in die Entwicklung von Prozessoren, die mit weniger Energie auskommen sollten, die Software-Produzenten entwarfen Betriebssysteme, welche die Leistung eines Notebooks automatisch und für den Nutzer unbemerkt drosseln, um so Energie zu sparen etc. Weil Nachfragesog-Innovationen über Marktsignale induziert werden, ist es wahrscheinlicher, dass diese sich eher durchsetzen werden als jene, die in Unkenntnis der Konsumentenwünsche und -akzeptanz auf vielversprechenden Patenten beruhen. Während die Entstehung von Verbesserungsinnovationen durchaus mit dem Nachfragesog-Modell plausibel erklärt werden kann, so unzulänglich erweist es sich, das Aufkommen von bahnbrechend neuen Technologien zu fokussieren.

Zur Erklärung von solch revolutionären Innovationen bietet sich das *Technologieschub-Modell* an. In Umkehrung des kausal-chronologischen Musters, das dem Nachfragesog-Modell zugrunde liegt, identifiziert dieses das *Angebot* von neuer Technologie als die treibende Kraft des gesellschaftlich-technologischen Wandels. Innovationen erweisen sich gegenüber bereits etablierten Produkten entweder in ihrer Funktionalität oder in ihrem Preis auf dem Markt als überlegen, so dass Konsumenten unweigerlich die Innovation präferieren. Das technologische Angebot schafft auf diese Weise seine eigene Nachfrage, ja sogar seinen eigenen Bedarf in der Gesellschaft. Diesem Modell liegt die Vorstellung zugrun-

de, dass es zu Innovationen aufgrund ihrer Markt- und technologischen Überlegenheit sowie des damit in Gang gesetzten Konsums keine Alternative gäbe. Die technologische Entwicklung wird demgemäß zum bestimmenden Faktor der technologisch-gesellschaftlichen Entwicklung. Für die Betonung der Angebotsseite spricht die Tatsache, dass radikale Innovationen sich durch zwei Mechanismen ihren Weg bahnen. Zum einen bietet ein Unternehmen eine radikale Innovation nicht nach dem Prinzip des »trial and error« an. Vielmehr investieren Unternehmen eine oft nur in Milliarden Euro bezifferbare Summe dafür, einen technologischen »Durchbruch« in die Richtung weiterzuentwickeln, die den antizipierten Konsumentenwünsche entgegenkommt. Viagra und orale Kontrazeptiva sind Beispiele für von Unternehmen erfolgreich praktizierte *push*-Strategien.

Zum anderen sind radikale Innovationen der Ausgangspunkt für ein ganzes Spektrum von Anwendungsmöglichkeiten, was die Wahrscheinlichkeit ungemein erhöht, dass aus dem Pool von Nutzungs-offerten eine einzige auf die Nachfrage von Konsumenten stoßen und somit zum Erfolg wird. Häufig ist damit solch ein Nutzungskonzept verbunden, das vom Erfinder selbst nicht intendiert wurde. Ein historisches Beispiel hierfür ist die Entwicklung des physikalischen Fernsprechapparates, der später als Telefon bekannt wurde (Rammert 1993a: 233f.). Die Erfinder glaubten ein Gerät erfunden zu haben, das der einseitigen Nachrichtenübermittlung (zur Weiterleitung von Anweisungen an Untergebene oder zur Übertragung von Konzerten) dient. Auch wenn diese Vision nur im kleinen Maßstab Wirklichkeit wurde, so motivierte das technische Prinzip der schnellen Nachrichtenfernübertragung eine Vielzahl von Ingenieuren dazu, ganz unterschiedliche Nutzungsvisionen durchzuspielen, bis letztendlich das Zwei-Kanal-Konzept, das eine Wechselrede ermöglichte, reüssierte und sich eine breite Nachfrage schuf. Eine ähnliche Geschichte ließe sich über den Computer erzählen. So prognostizierte der Leiter jenes Teams, das einen der ersten Computer konstruierte, für dieses Gerät eine Nachfrage von nicht mal ein halbes Dutzend Maschinen (Volti 1995: 44/45). Trotz der genannten Beispiele, die für einen Technologieschub sprechen, darf keineswegs davon ausgegangen werden, dass das Angebot einer radikalen technolo-

gischen Innovation automatisch eine Nachfrage erzeugt. Letztere entsteht häufig erst dann, wenn die Innovation in ein Netzwerk von Nutzern, investierenden Firmen, vorhandener Infrastruktur und politischem Protektorat integriert ist (Akrich 1992; Callon 1989; Garud/Karnøe 2003; Latour 1996). Wir kommen darauf später zurück.

Aufgrund der erwähnten Erklärungsmängel kann keiner der beiden Ansätze einen exklusiven Geltungsanspruch hinsichtlich der Entstehung von Innovationen beanspruchen. So ist nicht verwunderlich, dass es heute *state of the art* ist, *demand pull* und *technology push* bei der Analyse von Innovationsverläufen zu kombinieren und die linearen Sequenzen aus ›Angebot generiert Nachfrage‹, bzw. ›Nachfrage generiert Angebot‹ zu einem rekursiven Modell der Innovationsentstehung zu koppeln (Mowery/Rosenberg 1979) (vgl. Kap. III/2.2). Dabei wird differenzierend davon ausgegangen, dass der Technologieschub in den frühen Stadien der Innovation und der Nachfrageschub in den späteren an Bedeutung gewinnt (Walsh 1984).

1.2 Phasen – Deskription und nützliche ›Fiktion‹

Nachdem ich mit Nachfragesog und Technologieschub die beiden elementaren Mechanismen auf der Makroebene des gesellschaftlich-technologischen Wandels vorgestellt habe, wende ich mich der Mikroebene des organisierten Innovationsprozesses zu. Weil Industriebetriebe und staatliche Einrichtungen die Entstehung des Neuen nicht dem Zufall überlassen, sondern Innovationstätigkeiten mit Geld-, Sach- und Personalmittel forcieren, ist das Management an der Planbarkeit des eigentlich schwer Planbaren interessiert, um so Ressourcen rationell einsetzen zu können. Diesem Problem sich widmend, etablierte sich innerhalb der Betriebswirtschaftslehre (*micro economics*) das Fachgebiet »Innovationsmanagement« samt Lehrbüchern und -stühlen. Dort werden in einem Milieu zwischen Wissenschaft und Managementschulung die unterschiedlichen Phasenmodelle entwickelt und diskutiert, die je nach Interesse, Akzentsetzung und Detaillierung in einer beträchtlichen Anzahl im Umlauf sind (Verworn/Herstatt

2000). Häufig stützen sich diese Modelle auf Studien von erfolgreich durchgeführten Innovationsprojekten. Die dabei identifizierten Abläufe werden zum nachahmenswerten Vorbild stilisiert, inklusive der entsprechenden Handlungsempfehlungen. In solchen Fällen verschwimmt die Grenze zwischen deskriptivem und normativem Modell (Cooper 1983), was bei der Darstellung im Folgenden zu berücksichtigen ist. Das lineare Modell in seiner einfachen Ausprägung besteht aus vier idealtypischen Phasen (vgl. Tab. 2):

Tabelle 2: Phasen des linearen Innovationsprozesses

Idealtypische Phase des Innovationsprozesses	Institutionen
Entdeckung (<i>discovery</i>)	Institute und Laboratorien der Grundlagenforschung
Erfindung (<i>invention</i>)	Institute und Laboratorien der anwendungsorientierten Forschung
Entwicklung (<i>development</i>)	Forschungs- und Entwicklungsabteilungen in Unternehmen
Verbreitung (<i>diffusion</i>)	Nutzungskontexte der Innovation (Haushalte, industrielle Verwender etc.)

Das Stadium der *Entdeckung (discovery)* wird von Grundlagenforschung geprägt. Eine von ökonomischen Zwängen weitgehend befreite Forschung stößt auf solche Phänomene, die eine technologische Verwertung versprechen.

Im Stadium der *Erfindung (invention)* greift die anwendungsorientierte Forschung die von der Grundlagenforschung entdeckten Erkenntnisse auf, um sie in einen robusten und gezielt einsetzbaren Wirkungszusammenhang zu bringen. Der Prozess der Erfindung endet in der Regel mit der Patentierung.

Durch unzählige Veränderungen und Verbesserungen sind

Entwicklungsabteilungen von Unternehmen bemüht, aus der nun bekannten Technologie ein Produkt oder Verfahren zu entwickeln, das einem konkreten Nutzen dient. Im Stadium der *Entwicklung (development)* stehen ferner Fragen der Produzier- und Kommerzialisierbarkeit im Mittelpunkt.

Der Innovationsprozess endet mit der Phase der *Diffusion*. In dieser Phase entscheidet sich, ob die Innovation letztlich von Nutzern adaptiert oder abgelehnt wird.

Sämtliche linearen Phasenmodelle basieren auf der Annahme, dass der gesamte Innovationsprozess aus zeitlich aufeinander abfolgenden und distinkt unterscheidbaren Einheiten besteht. Beide Aspekte – der chronologische Ablauf und die Vorstellung von distinkten Phasen – werden von der neueren Forschung hinterfragt und relativiert:

- Wissenschaft und Forschung gelten heute nicht mehr als der alleinige oder alternativenlos wichtigste Auslöser von Innovationen (vgl. hierzu die Position von Bush 1998 [1945]). Empirische Studien weisen darauf hin, dass Erfindungsideen nicht nur am Beginn des Innovationsprozesses generiert werden, sondern auch in späteren Phasen der Entwicklung eine Rolle spielen. Sie können sich markant von den ursprünglichen Ideen unterscheiden und diese sogar an Wichtigkeit übertreffen (Du Gay et al. 1997; Garud/Karnøe 2003; Hughes 1976). Wissenschaft muss nicht der Ursprung des Innovationsprozesses sein, sondern häufig wird sie erst dann zu Rate gezogen, wenn es bei der Entwicklung und Verbreitung einer Innovation zu praktischen Problemen kommt (vgl. das Ketten-Modell unten in Kap. III/3.1).
- Das lineare Phasenmodell postuliert eine Gerichtetheit in der Abfolge des Innovationsprozesses. Dies widerspricht dann der Realität, wenn es zu Rückschlägen, Widerständen oder fortlaufenden Verbesserungen kommt. Neuere Innovationsmodelle greifen die Einsicht auf und zeigen, dass Abbrüche sowie Rückkopplungen zwischen den Phasen nicht nur die Ausnahme, sondern die Regel sind (Maidique/Zirger 1985).

Trotz seiner Kritisierbarkeit darf das Phasen-Modell nicht in

Bausch und Bogen verworfen werden. In Fällen, in denen neue Produkte für einen reifen, hochsegmentierten, spezialisierten Markt entwickelt werden, stellt dieses Modell durchaus eine Annäherung an den tatsächlichen Innovationsverlauf dar, insbesondere dann, wenn die Konsumenten nicht als Wissensquelle für die Entwicklung des neuen Produkts befragt werden, sondern nur als passive Käufer in Erscheinung treten (Fleck 2002: 15). Abschließend ist anzumerken: Als nützlich erweist sich das lineare Modell mehr in der betrieblichen als in der wissenschaftlichen Praxis. Die Unterscheidung in Phasen dient Organisationen dazu, Zäsuren für Entscheidungsprozesse zu setzen. Auf diese Weise entstehen Zeitsegmente, die den Entwicklungsprozess strukturieren und als Wegmarken Prozesskontrollen ermöglichen. Auf diese Weise werden Phasenmodelle zur ›nützlichen Fiktion‹ für das Management.

2. NON-LINEARE MODELLE

Die eben vorgestellten linearen Modelle zeichnen sich *erstens* dadurch aus, dass die Ursache für Innovationen eindeutig identifiziert ist: Der Konsum bei Nachfragesog-, der technologische Fortschritt bei Technologieschub-Modellen beziehungsweise Forschung und Entwicklung bei den Phasenmodellen auf der Mikroebene; *zweitens*, dass der technologische Wandlungsprozess als ein kontinuierliches Fortschreiten betrachtet wird; und *drittens*, dass dieses in einer Phasenfolge geschieht. Welche Beschreibungsalternativen gibt es? Diese Frage beantworten die nächsten Abschnitte.

2.1 Wellen, Kombinationen und Unternehmerscharen

Für eine Soziologie der Innovation ist das von Schumpeter entwickelte Modell grundlegend, weil ihm zufolge Innovationen nicht von Außen wie ein Himmelsgeschenk in die Gesellschaft treten, sondern als Produkte von gesellschaftlichen Prozessen beobachtet

werden. Diese Betrachtungsweise, die im Grunde ökonomische und soziologische Perspektiven integriert, schlägt sich in der Charakterisierung des Innovationsprozesses als Zyklus nieder, in dessen Verlauf sich Phasen der Innovation und der Imitation (das Kopieren von Innovationen durch Nachahmer) wechselseitig ablösen. Die Wirkung von Innovationen innerhalb der sozio-ökonomischen Entwicklung liegt in der Störung des wirtschaftlichen Gleichgewichts, die den Anfang einer ganzen konjunkturellen Wellenbewegung markiert. Innovationen sind Schumpeter zufolge die Ursache von zyklischen Wirtschaftsschwankungen, wenn das vermehrte Auftreten von Innovationen und deren Abklingen abwechselnd Prosperität und wirtschaftliche Depression hervorrufen. Um seine Theorie zu entfalten, verwendet Schumpeter einen ebenso abstrakten wie breiten Innovationsbegriff. Abstrakt ist sein Begriff deshalb, weil er an der für die Mikro- und Makroökonomie relevanten mathematisierten Produktionsfunktion ansetzt. »Wenn wir nicht die Faktormengen [also die Anzahl der Arbeitskräfte, des Rohmaterials etc.], sondern die Form der Funktion verändern, dann haben wir die Innovation«, heißt es lapidar (Schumpeter 1961: 94). Eine Innovation stellt demnach nichts anderes dar als die Infragestellung des bewährten Modus, den Ertrag einer Produktion zu berechnen. Breit ist der Innovationsbegriff von Schumpeter deshalb zu nennen, weil er sich nicht nur auf neue Produkte und Produktionsprozesse bezieht, sondern grundsätzlich jedwede unternehmerische »Durchsetzung neuer Kombinationen« umfasst (Schumpeter 1964 [1911]: 100). Diese Kombinationen lassen sich typisieren:

- die Herstellung von für Konsumenten neuartigen Produkten;
- die Einführung neuer Produktionsmethoden;
- die Erschließung neuer Absatzmärkte;
- das Aufspüren von neuen Bezugsquellen von Rohstoffen oder Halbfabrikaten;
- die Schaffung von neuen Organisationen (ebd.: 100/101).

In seiner Betrachtung des gesellschaftlich-wirtschaftlichen Wandels unterscheidet Schumpeter zwei Phasen: zum einen eine *Phase der Statik*, die durch mindere Anpassungsvorgänge inner-

halb eines existierenden Gleichgewichts gekennzeichnet ist; zum anderen eine *Phase der Perturbation* infolge innovatorischen Handelns, ehe sich nach einiger Zeit aufs Neue eine Gleichgewichtslage etabliert. Die erste Phase wird von der Figur des Unternehmers geprägt, die in Schumpeters Theorie eine Schlüsselposition einnimmt. Als Unternehmer gilt derjenige, dem es gelingt, neue Kombinationen durchzusetzen. Da Unternehmertum mit Innovativität unauflöslich verbunden ist, kann eine Person nur temporär eine solche Rolle übernehmen. »Niemand ist ununterbrochen Unternehmer, und niemand kann immer nur Unternehmer sein« (Schumpeter 1961: 111). Spätestens dann, wenn eine Person eine neue Kombination durchgesetzt hat, ist es an der Zeit, das Handlungsrepertoire des Unternehmers mit dem des Managers zu tauschen. Schumpeter grenzt die Rolle des Unternehmertums auf der einen Seite von der des Managers sowie Kapitalbesitzers ab und auf der anderen von der des Erfinders. Unternehmensein ist eine Funktion, die prinzipiell eine jede Person übernehmen kann, sofern sie über die notwendige Begabung verfügt (ebd.: 112). Um neue Kombinationen durchzusetzen, ist es nach Ansicht von Schumpeter nicht zwingend, Eigenkapital aufzubringen. Um als Unternehmerin zu fungieren, muss eine Person nicht notwendigerweise Betriebsbesitzerin sein, sie ist durchaus auch in der Position der Gehaltsempfängerin, wie z.B. Arbeiterin, Einkäuferin etc. anzutreffen. Daher ist die »Unternehmerfunktion« – wie Schumpeter schreibt – nicht ohne weiteres in einer Organisation zu lokalisieren, da sie mit den formalen Führungspositionen nicht prinzipiell übereinstimmen muss.

So wie für Schumpeter keine zwingende Verknüpfung von Unternehmertum und Fabrikbesitz besteht, so gibt es ebenfalls keine zwischen ersterem und dem Erfindertum. Historische Beispiele, wie Werner von Siemens oder Graham Bell, belegen für Schumpeter die Tatsache, dass die Erfindergabe und die Fähigkeit, unternehmerisch zu handeln, sich nicht wechselseitig ausschließen müssen. Doch wesentlich für die Unternehmerfunktion ist die Durchsetzung von neuen Kombinationen, nicht deren Erfindung. »The inventor produces ideas, the entrepreneur ›gets the things done‹, which may but need not embody anything that is *scientific*»

cally new« (Schumpeter 1991: 413). Schumpeter geht davon aus, dass fortlaufend neue Erfindungen dargeboten werden, doch es bedarf notwendig des Unternehmers, der diese als Neukombination erkennt oder mit Geschäftsideen kombiniert und in die Realität umzusetzen versteht (Schumpeter 1964 [1911]: 117). Der »schöpferische Unternehmer« ist einerseits derjenige, der über das Potenzial verfügt, Neuerungen zu verwirklichen und gegen Widerstände durchzusetzen, der andererseits jedoch als sozialer Aufsteiger nicht das Kapital besitzt, dieses Potenzial zu verwirklichen. Daher kommt den Financiers aus dem Bankwesen neben den Unternehmern eine Schlüsselrolle zu, weil sie durch Kreditvergabe den Pionier erst handlungsfähig machen (Schumpeter 1961: 119).

Nachdem der Innovator – die Hemmnisse überwindend – eine Neuerung etabliert hat, lockt sein Erfolg Nachahmer an, welche die Innovation nicht nur aufgreifen, sondern auch durch Verbesserungen häufig, was den Ertrag oder die Effizienz angeht, übertreffen. In dieser zweiten Phase des Innovationsprozesses, die vom Imitationsverhalten geprägt ist, treten Unternehmer scharenweise auf, was Schumpeter (1964 [1911]: 339f.) als das Anzeichen eines beginnenden Aufschwungs wertet. Die Häufung von Unternehmensgründungen wird durch zwei Faktoren begünstigt: Zum einen dadurch, dass Nachfolger eine bereits durch Erfolg ausgewiesene Idee kopieren können und zum anderen durch die Kreditvergabepraxis der Banken. Nacheiferer erhalten für ihre Projekte deswegen weniger zögerlich Kredite als vormals die Pioniere, weil sich die Innovation bereits als tragfähiges Geschäftsmodell erwiesen hat. Neben den Imitationseffekten sorgen weitere Investitionen zur Verbesserung und Ausweitung der Innovation und damit für einen konjunkturellen Aufschwung. Der gleichzeitige Anstieg der Güterproduktion und der Preise führt unweigerlich zum Boom. Als bald jedoch schmälern sich die Unternehmensgewinne, weil im Vergleich zur Anfangszeit, als die Innovation noch jung war, günstigere Produkte den Markt überschwemmen. Dies führt dazu, dass der Wettbewerb nun über den Preis und nicht mehr über das Merkmal »Innovativität« ausgetragen wird. Die Preise fallen, die Innovationstätigkeit sinkt, ein Se-

lektionsprozess unter den Wettbewerbern tritt ein. Am Ende dieser Rezession pendelt sich wieder ein Gleichgewicht ein, das auf den im Innovationsgeschehen etablierten Strukturen beruht.

Schumpeters Entwicklungstheorie gilt als wichtiger Wegbereiter einer Evolutionstheorie der Innovation. Die von Karl Marx übernommene Idee Schumpeters, eine Volkswirtschaft anstatt ausschließlich in den Begriffen des Gleichgewichts in solchen der gestörten Balance (z.B. »schöpferische Zerstörung«) zu analysieren, ist für die Evolutionstheorie der Ausgangspunkt ihrer Analysen (Clarke/Juma 1988: 211f.; Elliot 1980).

2.2 Paradigmen, Trajektorien und Nischen

Ohne Übertreibung lässt sich behaupten, dass der evolutionäre Ansatz in der Tradition Schumpeters der heutigen Innovationsforschung als einendes Basiskonzept dient. Allerdings ist es im gegenwärtigen Stadium der Diskussion immer noch verfrüht, von einem kompletten und kohärenten theoretischen System auszugehen (Saviotti/Metcalf 1991: 2). Vielmehr ist zu beobachten, dass einzelne Autorinnen und Autoren den Begriffsapparat der biologischen Evolutionstheorie mal metaphorisch, mal als Werkzeugkasten verwenden, um bestimmte Aspekte ihrer Überlegungen zu plausibilisieren und zu veranschaulichen. Bei der Darstellung der Grundrisse der evolutionären Theorie von Innovationsprozessen konzentriere ich mich auf die grundlegende – auch von Schumpeter übernommene – Idee, dass sich der technologisch-gesellschaftliche Wandel in zwei Phasen unterteilen lässt (Dosi 1982; Dosi et al. 1988; Mensch 1975; Sahal 1985; Tushman/Anderson 1986): Eine *Phase inkrementeller Innovationen*, die durch eine *Phase radikaler bzw. revolutionärer Innovationen* unterbrochen wird. Um diesen Phasenwechsel nachzuzeichnen, beschränke ich mich weitgehend auf die Studien von Giovanni Dosi, weil es ihm gelungen ist, die Entstehungsgründe für die beiden eben benannten Innovationsphasen zu erklären. Dabei stützt er sich auf die Begriffe des wissenschaftlichen Paradigmas bzw. Paradigmawechsels des Wissenschaftshistorikers Thomas Kuhn (1976), um diese auf das Feld technologischer Entwicklungen zu übertragen.

Des Weiteren steuert Dosi Überlegungen bei, wie die bereits diskutierten klassischen Technologieschub- und Nachfragesog-Modelle (siehe Kap. III/1.1) miteinander verknüpft werden können (Dosi 1982: 151f.).

In dieser Perspektive vollzieht sich der technologische Fortschritt während einer ersten Phase in inkrementellen Schritten. Unternehmen greifen bei der Entwicklung von Technologien auf bewährte Suchroutinen zurück (Nelson/Winter 1977). Derlei Innovationen folgen einem Nachfragesog (*demand pull*), da die Unternehmen diejenigen Eigenschaften weiterzuentwickeln trachten, die den durch Marktstudien eruierten Kundenwünschen entsprechen. Bestehende Technologien, die sich bereits auf dem Markt bewährt haben, werden zum Ausgangspunkt der innovatorischen Bemühungen. Hierin liegt das konservative Moment der Innovationstätigkeit. So bestehen die Innovationen im Automobilantrieb weitestgehend darin, den Benzinmotor sparsamer, billiger, leichter, leiser, leistungsfähiger, pannenresistenter zu machen; weniger darin, einem ganz neuen Antriebsstoff (wie zum Beispiel Wasserstoff) zum Durchbruch zu verhelfen. Letztlich geht es darum, das Bestehende zu optimieren, weshalb diese inkrementellen Fortschritte auch Verbesserungsinnovationen genannt werden (Mensch 1972). Diejenigen Entwicklungs- und Suchroutinen, die in einer gewissen Erwartbarkeit zu inkrementellen Innovationen führen, bezeichnet Dosi als »technologisches Paradigma«. Ein solches umfasst Musterlösungen für eine bestimmte Auswahl von technologischen Problemen und basiert auf einer Auswahl von naturwissenschaftlichen Prinzipien und verwendeten Materialien (Dosi 1982: 152). Dadurch, dass mit dem Paradigma festgelegt wird, welche technologischen Probleme relevant und welche Lösungsmethoden opportun sind, bestimmt dieses die Richtung zukünftiger Innovationen. Es zeigt dem Management außerdem an, wo die Forschungs- und Entwicklungsinvestitionen zu akkumulieren sind. Damit gibt das Paradigma Ingenieur/-innen und Manager/-innen eine Orientierung und blendet zugleich alternative technologische Möglichkeiten aus.

Die Verwendung des von Kuhn entliehenen Paradigmbegriffs impliziert zwei Folgeüberlegungen: *Erstens* leitet das Paradigma nicht nur die Forschungs- und Entwicklungsbemühungen von

Ingenieur/-innen an, sondern es fungiert gleichzeitig als vergemeinschaftendes Band. Das heißt, die Anhängerschaft eines Paradigmas bedingt gleichzeitig die Zugehörigkeit zu einer technologischen Gemeinschaft (hierzu mehr in Kapitel IV/2). Zweitens geht, wie technikhistorische Studien belegen, ein neues technologisches Paradigma nicht linear aus dem alten hervor (Constant 1980; Geels 2002a). Verbesserungsinnovationen führen in der Regel nicht zu einer gänzlich neuen Technologie. Um die Kontinuität des inkrementellen Fortschritts zu brechen, bedarf es eines Paradigmawechsels (Dosi 1988a: 228), der gleichzeitig die Gemeinschaft und die Unternehmen, die ihren Erfolg dem bewährten Paradigma verdanken, herausfordert und letztlich ablöst. Solche radikalen Innovationen, wie sie Dosi nennt, manifestieren in besonderem Maße den gesellschaftlichen Aspekt des Innovationsprozesses, weil sie mit dem Auf- bzw. Abstieg von technologischen Gemeinschaften und Unternehmen verbunden sind.

Wie kommt es jedoch zu einer radikalen Innovation? Was sind die Auslöser? Radikale Innovationen können durch zwei Konstellationen ausgelöst werden:

Erstens: Verwender nehmen bei einem dominierenden Paradigma mehr und mehr *funktionale Mängel* wahr. Selbst wenn eine Technologie in ihrer Geschichte fortwährend verbessert und ihre Anwendungsweisen vermehrt werden, bleiben Probleme bestehen, die lange Zeit nicht als solche wahrgenommen werden müssen. Dass ein herkömmlicher Automobilmotor Benzin verbrennt, wird solange nicht als Defizit erachtet, bis der Rohstoff zu einem kritischen Kostenfaktor beim Gebrauch wird. Je mehr sich Benzin verteuert, desto mehr fällt der Benzinverbrauch als das ins Gewicht, was Constant (1984: 30f.) einen funktionalen Mangel nennt. Erst mit der Wahrnehmung dessen wird die Suche nach Alternativen motiviert. Solche funktionalen Mängel ergeben sich nicht nur, wenn sich das sozio-ökonomische Umfeld einer Technologie verändert, sondern auch dann, wenn sich ein ganzes System aus miteinander verknüpften Einzeltechnologien weiterentwickelt und eine einzelne im Kontext dieser Entwicklung als rückständig erlebt wird (Hughes 1987: 73). Zum Beispiel galt die analoge Fotografie, was die Aufnahmequalität anbetrifft, lange Zeit als der digitalen überlegen. Doch dieser Vorzug verlor an Be-

deutung, je mehr Nutzer über einen Computer verfügten, mit dem sie digitalisierte Bilder ansehen, bearbeiten, ausdrucken und verschicken konnten. Zwar mag die Analogphotographie bis heute *in puncto* Bildqualität der digitalen ebenbürtig sein, doch sie lässt sich nicht ohne weiteres in das digitalisierte System aus Computer und Internet integrieren. Sie wird als rückständig erlebt, weil sich der Kontext ihrer Anwendung verändert hat.

Zweitens: Neben den eben erwähnten funktionalen Mängeln können *wissenschaftliche Fortschritte* ein bestehendes technologisches Paradigma in einem anderen Licht erscheinen lassen. So mag das Paradigma noch Erfolge zeitigen und nach wie vor ein beträchtliches Entwicklungspotenzial bieten, doch die Befunde aus der wissenschaftlichen Forschung lassen Eingeweihte errahnen, dass zukünftige Entwicklungen im entsprechenden Feld auf anderen natur- und technikwissenschaftlichen Prinzipien beruhen könnten. So spielte zum Beispiel in der Flugzeugentwicklung die Aerodynamik die Rolle derjenigen wissenschaftlichen Spezialdisziplin, welche die Entwicklungsroutinen des Propellerantriebes zu hinterfragen begann und den Weg für den Düsenantrieb ebnete. So kam bereits Ende der 1920er Jahre die aerodynamische Forschung zu dem Ergebnis, dass es strahlgetriebenen Fluggeräten möglich ist, sich bis an die Grenze der Schallgeschwindigkeit zu beschleunigen, wohingegen Propellerantriebe hierfür nicht geeignet sind (Constant 1990: 226).

Anders als bei Verbesserungsinnovationen, die durch Marktsignale ausgelöst werden können, spielen für Dosi wissenschaftliche Fortschritte eine entscheidende Rolle, technologische Möglichkeiten zu eröffnen. Die Entstehung von bedeutenden technologischen Paradigmen ist oftmals abhängig von größeren naturwissenschaftlichen Durchbrüchen (Dosi 1984). Man denke hierbei zum Beispiel an die Biotechnologie oder an die synthetische Chemie. Gerade in der Entstehungsphase von neuen technologischen Paradigmen kommt es zu einer Konvergenz von wissenschaftlichen und ingenieurtechnischen Interessen und Praktiken (Dosi 1988b: 123) und zur Vermischung von ingenieur- und naturwissenschaftlichen *Communities*. Dabei beeinflusst die Forschung nicht nur die technologische Entwicklung, sondern die Erfindung von technologischen Instrumenten (man denke z.B. an

das Elektronenrastermikroskop oder an den Computer) trägt wesentlich zu den Fortschritten in den Naturwissenschaften bei. Darüber hinaus ist zu beobachten, dass speziell in den frühen Phasen der Paradigmenentwicklung technologische Problemlösung und theoretischer Erkenntnisgewinn nahe beieinander liegen. So wird die ingenieurtechnische Entwicklung von Geräten, chemischen Verbindungen etc. oftmals als Test für den Wahrheitsgehalt einer naturwissenschaftlichen Theorie betrachtet. Die ingenieur- und naturwissenschaftlichen Gemeinschaften gehen dann wieder getrennte Wege, sobald das technologische Paradigma etabliert ist. Wenn eine wissenschaftlich-technologische *Community* ein neues technologisches Paradigma auf den Weg bringt, dann folgen eine Reihe zeitaufwendiger und kostenintensiver Weiterentwicklungen durch »überbrückende Institutionen« (Dosi 1982: 155) wie z.B. öffentlich oder privat finanzierte Institute der angewandten Forschung. Um die Lücke zwischen der »Entdeckung« eines technologischen Paradigmas und dem Produkt, das zu vertretbaren Kosten hergestellt werden kann, zu schließen, sind unzählige unspektakuläre Probleme zu lösen, wobei man sich niemals sicher sein kann, wie viel Aufwand es zu deren Überwindung bedarf (Rammert 2002: 177). Ein großer Anteil dieser Entwicklungsarbeit ist auf das so genannte »*scaling up*« gerichtet. Es ist eine Sache, ein technologisches Artefakt (z.B. ein Medikament) im Labor zu entwickeln, eine andere ist es, dieses in Massenproduktion herstellen zu können (Volti 1995: 37). Es ist eine Sache, einen technologischen Prozess innerhalb des Labors zum Funktionieren zu bringen, eine andere, diese Funktion auch in »natürlicher« Umgebung zu gewährleisten, wo Störfaktoren nicht mehr oder nur unzureichend zu kontrollieren sind (Latour 1983).

Eine wichtige Rolle bei der weiteren Entwicklung eines technologischen Paradigmas spielen Nischen. Es handelt sich hierbei um einen geschützten Bereich, in dem eine Innovation von den selektierenden Effekten des freien Marktes geschützt ist (Rip/Schot 2002: 16ff.). Auch wenn man dazu eine kritische Haltung einnimmt, muss man konstatieren, dass das Militär eine klassische Nische darstellt, in der neue technologische Paradigmen weiterentwickelt und getestet werden (Beispiele sind das Internet,

das »Global Positioning System« oder die Kernspaltung). Nischen entstehen auch durch politische Regulierungen. So kreierte die dänische Regierung durch öffentliche Förderung von Windenergie insofern einen Nischenmarkt, als es infolge dieser politischen Intervention für viele Farmer/-innen lukrativ wurde, Windturbinen auf dem eigenen Grundstück zu betreiben oder betreiben zu lassen (Karnøe 1999). Solche durch politische Interventionen geschaffenen Nischen sind häufig Vorläufer für Marktnischen und tragen aus verschiedenen Gründen zum »take off« eines technologischen Paradigmas bei (Kemp et al. 2001: 275): *Erstens* entstehen in Nischen »Demonstrationsobjekte« (*hopeful monsters*), die bei Investoren Interesse wecken können; *zweitens* wird in Nischen ein interaktiver Lernprozess angestoßen, der dazu führt, dass komplementäre Techniken, die zur Funktionalität des technologischen Paradigmas beitragen, entwickelt werden. Eine komplementäre Technik zum technologischen Paradigma der Internetkommunikation ist das »http-Protokoll«, das es den Nutzern erst möglich macht, durch die Netze zu »surfen«. An diesem Lernprozess ist auch die Politik beteiligt, wenn sie durch Gesetzgebung und durch finanzielle Zuwendungen die Bedingungen in der Nische regelt und auch dafür sorgt, dass dieser geschützte Bereich zu einem geeigneten Zeitpunkt den Marktkräften anheim gegeben wird.

Spätestens dann, wenn ein neues technologisches Paradigma jenseits der Nische auf Interessenten stößt, beginnt innerhalb der Phase der radikalen Innovation das Stadium der Fermentierung, wenn das alte mit dem neuen Paradigma rivalisiert (Tushman/Rosenkopf 1992: 318f.). Man spricht deshalb von einem »Gärungsstadium«, weil es nicht ausgemacht ist, dass sich das neue technologische Paradigma zwingend durchsetzt. Ein Beispiel für eine solche Konkurrenz zweier Paradigmen ist der Wettbewerb zwischen dem *Open Source*-Betriebssystem »Linux« und dem etablierten, von »Microsoft« entwickelten »Windows«. Das dominierende Paradigma hat allein dadurch, dass es bereits adaptiert ist, Vorteile auf seiner Seite: Zum einen sind die Nutzer/-innen mit der Technik vertraut. Eine Umstellung auf »Linux« würde für Betriebe und Verwaltungen Investitionen in die Fortbildung der Mitarbeiter/-innen mit sich bringen. Zum anderen sind komple-

mentäre Techniken auf das dominierende Paradigma abgestimmt. Jenseits eines direkten Funktionsvergleiches zwischen den rivalisierenden Betriebssystemen liegen die Vorzüge von »Windows« gegenüber »Linux« darin, dass die bereits angeschaffte Anwendungssoftware und die darin gespeicherten Datensätze als »Windows«-kompatibel gelten. Für eine Umstellung auf das neue Paradigma werden an etlichen Stellen Kompatibilitätsprobleme befürchtet. Eine dominante Technologie wird durch ein ganzes *Regime* stabilisiert, welches neben wissenschaftlichem Wissen und ingenieurwissenschaftlichen Praktiken insbesondere auch die bereits vorhandene Infrastruktur und die Verwendungsroutinen umfasst (Kemp et al. 2001: 272f.; van de Poel 2003: 50). Während des Fermentierungsstadiums ist zu beobachten, wie sowohl das neue als auch das alte Regime mit Produktvariationen aufwarten (Tushman/Rosenkopf 1992: 318). Häufig reagiert das dominante Regime auf die Herausforderung eines neuen Konkurrenten mit einer gesteigerten Innovativität. Es kommt zu einem »Segelschiff-Effekt« (Geels 2002b: 379), dessen Namen auf die historische Begebenheit anspielt, dass Segelschiffe durch die Herausforderung der Dampfschiffahrt enorm verbessert wurden, was Segel-, Rumpfgroße und Geschwindigkeit anbetrifft. Durch diesen Entwicklungsschub des alten Regimes konnte es eine Zeit lang neben dem neuen koexistieren. Das Fermentierungsstadium ist durch eine Ergebnisoffenheit gekennzeichnet, weil nicht vorherbestimmt ist, ob sich das Neue und vermeintliche Bessere tatsächlich durchsetzen wird.

Der evolutionäre Mechanismus, der darüber entscheidet, ob sich ein technologischer Paradigma- und damit auch Regimewechsel vollzieht, ist die Selektion durch die Umwelt in Gestalt des Marktes (Metcalfe 1988: 568f.; Nelson/Winter 1977: 64f.). Aufgrund ihrer Überlegenheit gegenüber bewährten Gütern erzeugen Produktinnovationen eine erhöhte Nachfrage beim Konsumenten. Prozessinnovationen verbilligen in der Regel die Produktion. Beides führt dazu, dass das innovierende Unternehmen seine Gewinnmargen erhöht. Damit erweitert sich dessen Spielraum für Neuinvestitionen oder zur Verbilligung seiner Produkte, wodurch sich die Gewinnchancen für nicht innovierende Konkurrenten schmälern und diese auf lange Sicht vom Markt gedrängt

werden (Dosi 1988b: 145). Doch Märkte dürfen keinesfalls als Instanzen betrachtet werden, die beobachterunabhängig, interessenneutral und wissenschaftlich objektiv entscheiden, welche Technologie die bessere ist. Märkte sind – gerade dies ist der unhintergehbare Standpunkt einer soziologischen Betrachtung (vgl. Callon 1998; Fligstein 1996; White 1981) – gesellschaftliche Konstruktionen, sie werden mitkonstituiert durch politische Interessen, kulturelle Wandlungen, rechtliche Vorgaben, internationale Standards und institutionelle Mechanismen wie z.B. Vertrauen. Der Markt in Gestalt des Aufeinandertreffens unabhängiger, rational entscheidender und über vollständige Marktübersicht verfügende Akteure ist aus mehreren Gründen eine idealtypische und oftmals realitätsferne Vorstellung. Zunächst werden Märkte – wie obige Ausführungen über Nischen bereits darlegen – für Innovationen als selektive Umwelt durch politische Interventionen geschaffen.

Selbst Massenmärkte sind durch die Setzung technischer Standards oder Normen politisch beeinflusst (Jørgensen/Sørensen 1999). So verändert zum Beispiel die Verschärfung der Emissionsnormen die Marktbedingungen im Automobilgewerbe in der Weise, dass Innovationen in eine bestimmte Richtung gelenkt werden (Schot/Rip 1996: 258f.). Ferner ist es inzwischen *common sense*, dass das Entscheidungsverhalten von Marktteilnehmern durch deren »bounded rationality« (Simon 1955, 1997) begrenzt ist. Das heißt, die Frage, ob eine Innovation dem Bewährten überlegen ist, werden etwaige Käufer/-innen nicht durch einen unendlichen Klärungsprozess, sondern durch soziales Verhalten zu beantworten suchen. Anstatt »rational« zu entscheiden, neigen Käufer/-innen häufig zur Imitation, wenn sie Innovationen deshalb annehmen, weil es andere schon vor ihnen getan haben. Nicht zu vergessen ist, dass die Präferenz für das Innovative in der Moderne mit dem Erwerb von Prestige verbunden ist. Insbesondere die Soziologie hat darauf aufmerksam gemacht, dass nicht die Umwelt über die Evolution von Innovationen in einer »trial and error«-Logik entscheidet, sondern dass die Umwelt in Gestalt von Nutzern durchaus auch auf die Struktur der Innovation Einfluss nimmt und nicht bloß »blind« selektiert. Dies geschieht zum Beispiel in sogenannten Nutzer-Produzenten-Interaktionen. Speziell

bei solchen Produkten, bei denen Nutzer sich einen großen Vorteil von einer Innovation versprechen (wie z.B. Forscher/-innen von wissenschaftlichen Messgeräten), werden diese, wenn ihre spezifischen Anforderungen und Erwartungen an den Produzenten adressiert werden, zur Quelle für weitere Innovationen (von Hippel 1976). Die Grenze zwischen der selektierenden Umwelt des Marktes und des Unternehmens wird durch derartige Interaktionen zu einem gewissen Grade aufgelöst (zur Gestaltung von Märkten durch Unternehmen vgl. Rupp 1999).

Bewährt sich ein neues technologisches Paradigma in seiner Umwelt, ist die nun entstandene radikale Innovation wiederum die Basis für den ›normalen‹ kontinuierlichen technischen Wandel in Form von unzähligen inkrementellen Verbesserungen und Modifikationen. Es ist damit ein Pfad für den weiteren Innovationsverlauf vorgegeben. Solche auf einem technologischen Paradigma basierenden inkrementellen Weiterentwicklungen beschreiben eine *Trajektorie* (Verlaufsbahn) (Dosi 1982: 153, 1988b: 115f.; Nelson/Winter 1977: 56f.). Durch eine Diffusion des technologischen Paradigmas und vor allen Dingen durch weitere Verbesserungen und solche Innovationen, welche die Funktionalität des technologischen Paradigmas steigern, stabilisiert sich eine Trajektorie, wodurch der Paradigma- bzw. Regimewechsel irreversibel wird (Hughes 1987: 76f.).

Welche evolutionären Stabilisierungsmechanismen sind im Einzelnen zu beobachten?

1. *Learning by using*: Je häufiger eine Technologie benutzt wird, desto mehr Erfahrungswissen wird produziert. Dies erleichtert den Umgang mit ihr, so dass ein Wechsel zu Konkurrenzprodukten, über die man nicht Bescheid weiß, weniger wahrscheinlich wird (Arthur 1988).
2. *Netzwerkexternalitäten*: In vielen Fällen steigt die Attraktivität einer Technologie, je mehr Nutzer daran teilhaben (ebd.). So steigt zum Beispiel der Nutzen der E-Mail-Kommunikation, je mehr Personen auf einen E-Mail-Account zurückgreifen können.
3. *Technologische Interrelationen*: Eine Technologie wird umso mehr von Nutzern adaptiert, je mehr unterstützende Techno-

logien Teil der Infrastruktur werden (ebd.). Die Überlegenheit des benzinbetriebenen Automobils besteht darin, dass ein flächendeckendes Netz von Tankstellen, Autoreparaturwerkstätten etc. den Umgang komfortabel macht, was eine alternative Antriebstechnologie, wie zum Beispiel der Wasserstoffmotor, der in keine derartige Infrastruktur eingebettet ist, nicht ohne weiteres anbieten kann.

All diese stabilisierenden Effekte sorgen dafür, dass ein einmal eingeschlagener technologischer Pfad nicht ohne Widerstände verlassen werden kann. Es entsteht in der Weise eine Pfadabhängigkeit (vgl. IV/2.1). Der gesellschaftlich-technologische Wandel wird *pfadabhängig*, weil die Innovationstätigkeit in der Gegenwart von zeitlich weit zurückliegenden Ereignissen bestimmt wird (David 1985). Eine solche Pfadabhängigkeit kann dazu führen, dass eine Technologie weiterentwickelt und benutzt wird, obgleich es in der Zwischenzeit vermeintlich bessere Alternativen gibt. All die oben genannten Stabilisierungsmechanismen stützen ein etabliertes technologisches Regime. Die technologisch-gesellschaftliche Entwicklung kann auf diese Weise in einem inferioren technologischen Paradigma eingeschlossen sein. Man spricht von dann von einem »lock-in« (Arthur 1989). Das meistzitierte Beispiel hierfür ist die QWERTY-Tastaturanordnung, die keineswegs die bequemste und schnellste Schreibweise ermöglicht.

3. METAPHERN DES INNOVATIONSPROZESSES

Auf der Ebene des gesellschaftlich-technologischen Wandels dient der evolutionäre Ansatz als vereinheitlichender Theorierahmen für die unterschiedlichen Beiträge der Forschung. Weniger übersichtlich erscheint die Theorielandschaft auf der Mikroebene des organisierten Innovationsprozesses. Kein bestimmendes Paradigma ist in Sicht. Lediglich zwei Merkmale lassen sich aufführen, die sämtliche Theorieangebote auf dieser Ebene gemeinsam haben: Sie artikulieren eine explizite Abkehr vom linearen Phasenmodell und greifen auf Metaphern zurück, um die Prozesslogik zu charakterisieren. Die Kette, das Rugby-Spiel und das Feu-

erwerk bilden die metaphorischen Bezugspunkte jener drei Theorieofferten, die ich vorstellen möchte. Meine Auswahl entzieht sich in der Weise des Vorwurfs der Beliebigkeit, als ich mich in meiner Darstellung an Ansätze halte, die in der Innovationsforschung jenseits allen notwendigen Kritikbedarfs Reputation erworben haben. Ein Anspruch auf Vollständigkeit oder Systematik wird nicht erhoben.

3.1 Ketten-Modell: Feedbackloops und Forschung als permanente Ressource

In Ermangelung eines Konzepts, das die Ergebnisse der damalig aktuellen Innovationsforschung integrierte, entwarfen Stephen Kline und Nathan Rosenberg (1986) das Ketten-Modell (*chain-linked model*). Bis heute gehört es zum Standardwissen der Disziplin (Hage/Hollingsworth 2000; Palmberg et al. 1999; Senker/Faulkner 1996). Das Modell findet seinen Ausgangspunkt in der Feststellung, dass dem Innovationsprozess grundsätzlich Unsicherheit innewohnt. Diese wird deshalb zum Problem, weil sie sich mit den herkömmlichen Methoden der Messbarkeit und Risikoabschätzung kaum absorbieren lässt. Je radikaler die Innovation, desto mehr steigt die Unsicherheit hinsichtlich der Mach-, Finanzier- und Durchsetzbarkeit derselben (Kline/Rosenberg 1986: 294). Organisationen sind einer solchen Situation jedoch nicht hilflos ausgeliefert, sondern haben jene Bewältigungsstrategien und -taktiken entwickelt, die im Ketten-Modell festgehalten sind. Explizit setzen sich Kline und Rosenberg vom linearen Phasenmodell (vgl. Kap. III/1.2) ab. Dieses repräsentiert ihrer Ansicht nach den Innovationsprozess als gleitenden, gut beherrschbaren, messbaren und planbaren Vorgang, wodurch die technologischen, organisatorischen und ökonomischen Unsicherheiten verborgen werden. Obendrein wird der Innovationsprozess als singulärer Vorgang beschrieben, dessen Ursprung in der Forschungs- und Entwicklungsabteilung verortet wird. Drei Befunde der Innovationsforschung lässt das lineare Modell nach Ansicht von Kline und Rosenberg außer Acht:

- Es verzeichnet keine *Feedback*-Schleifen zwischen den Phasen.
- Fehler und Fehlschläge als nicht substituierbare Elemente des innovatorischen Lernprozesses werden ignoriert.
- Das lineare Modell schätzt die Rolle von Wissenschaft als initiiertes Instrument im Innovationsprozess falsch ein.

Das Eingangstor für Innovationen liegt für Kline und Rosenberg vielmehr im Weiterentwicklungsprozess. Um dies deutlich zu machen, prägen sie den Begriff »analytisches Design« (*analytic design*) (ebd.: 302). Dies umfasst die Suche nach neuen Kombinationen bereits existierender Produkte und die Entwicklung von neuen Produkten innerhalb des Rahmens eines vorgefundenen Standes der Technik. Als Alternative zum linearen Phasenmodell schlagen Kline und Rosenberg ihr Ketten-Modell vor, das aus insgesamt fünf Prozesslinien besteht, wobei ich mich auf die drei wichtigsten in meiner Darstellung beschränke:

Die *erste* zentrale Kette der Innovation unterscheidet sich nicht vom linearen Modell. Sie beginnt mit der Erfindung und endet mit dem Marketing. Die *zweite* besteht aus der Reihe der Feedbackschleifen, die sich zwischen den Segmenten ereignen. Diese Rückkopplungsschleifen sind Ausdruck einer ständigen Koordination zwischen den Fachabteilungen. Das für den Innovationsprozess entscheidende Feedback stammt von der Beobachtung der Marktsignale bzw. der Kundenbedürfnisse, weil diese sogleich zur Verbesserung des Produktes und des Services verwertet werden (Barnekow 2004). Diese enge Rückkopplung von Marktbeobachtung und technischer Entwicklung ist für Kline und Rosenberg Motiv, die Trennung von Nachfragesog und Technologieschub (vgl. Kap. III/I.1) abzulehnen. Denn jedes Konsumentenbedürfnis, das von der Marktforschung erfasst wird, führt früher oder später zu einem neuen Design, und jede erfolgreiche und akzeptierte Innovation führt unweigerlich zur Veränderung der Marktverhältnisse (ebd.: 290).

Mit der *dritten* Kette wird die Relation von Wissenschaft zu den anderen Kettengliedern bestimmt. In Abgrenzung zum linearen Modell gehen Kline und Rosenberg davon aus, dass sich die Rolle von Wissenschaft nicht allein am Anfang eines Innovationsprozesses entfaltet, sondern sich entlang der ganzen zentralen ersten

Kette erstreckt. Das heißt, Wissenschaft ist mit jedem anderen Kettenglied verbunden. Solche Verknüpfungen können Kline und Rosenberg nur deswegen identifizieren, weil sie die Rolle von Wissenschaft im Innovationsprozess neu bestimmen. Wissenschaftliches Wissen wird darin in doppelter Weise relevant: *zum einen* in Form des verkörperten Wissens, das die Mitarbeiter/-innen als Ergebnis eines wissenschaftlichen Ausbildungsweges und einer Hochschulsozialisation in das Unternehmen einbringen. Zu dieser Form des Wissens gehört auch die Fähigkeit der Mitarbeiter/-innen, nach den benötigten Informationen zu recherchieren. *Zum anderen* ist jenes wissenschaftliche Wissen relevant, das bei Bedarf in den Laboratorien gewonnen wird und in sämtlichen Phasen des Innovationsprozesses angefordert werden kann und muss. Unersetzbar ist demnach die Rolle von Forschung – nur eben nicht als initiierender Funke, sondern als permanente Wissensressource entlang der ganzen Prozesskette. Kostspielige und zeitaufwendige Forschung wird erst dann eingesetzt, wenn diese eben genannte Wissensform an ihre Grenzen stößt.

Das Ketten-Modell stellt eine Erweiterung des Phasenmodells dar, wobei die unterschiedlichen Stadien durch Rekursionsschleifen verbunden sind. Zwar gelten solche Rekursionen heute als ein wichtiges Kennzeichen von Innovationsprozessen. Seit den 1990er Jahren zeichnet sich allerdings ab, dass die Phasen, die im Ketten-Modell nach wie vor strukturprägend sind, mehr und mehr durch Synchronisierungen aufgelöst werden. Einzelbeiträge eines Innovationsprojektes, die ehemals zeitlich sequenzialisiert wurden, finden nun nahezu zeitgleich statt. Außerdem ist kritisch anzumerken, dass das Ketten-Modell nicht sämtliche für den Innovationsprozess wichtigen Rückkopplungen erfasst. Man denke hierbei an avancierte Nutzer/-innen (so genannte Pilotkunden oder *leading edge user*), die in den Entwicklungsprozess integriert werden (Iansiti/MacCormack 1997; Rupp 1999: 371), oder all die Rückkopplungen, die zu beobachten sind, wenn mehrere Unternehmen an einem Innovationsprojekt beteiligt sind (vgl. Kap. IV).

3.2 Rugby-Modell: Überlappungen und Wissensteilhabe (*knowledge sharing*)

Werden im Ketten-Modell die Phaseneinheiten rekursiv verknüpft, so bleiben doch die zeitlichen Sequenzen erhalten. Im Vergleich dazu werden im Rugbymodell des Innovationsprozesses, das Hirotaka Takeuchi und Ikujiro Nonaka (1986) zur Diskussion stellen, die einzelnen Zeitphasen miteinander verschmolzen. Mitarbeiter/-innen, die ansonsten über den Zeitverlauf verteilt nacheinander zum Einsatz kommen, arbeiten hier in einem Team zusammen. Wie der Ball zwischen den Spielern in der Weise hin und her, vor und zurück gepasst wird, so dass sich eine Mannschaft als Einheit über das Feld bewegt, so soll das Team Teilprobleme des Projektes in Selbstorganisation an die jeweiligen dafür spezialisierten Teilnehmer/-innen delegieren. Um die Besonderheiten des Rugby-Modells hervortreten zu lassen, kontrastieren Takeuchi und Nonaka dieses mit dem linearen Phasenmodell (vgl. Kap. III/I.2), das sie als eine Art Staffellauf charakterisieren. In einem solchen reicht eine Spezialistengruppe den Staffelstab – in Form eines Teilprojektergebnisses – an die nächste Arbeitsgruppe weiter. Auf diese Weise schreitet das Projekt von Sequenz zu Sequenz voran. Bei dieser Prozessorganisation nutzt das Modell die Vorteile der Arbeitsteilung und der Spezialisierung der Mitarbeiter/-innen aus, indem die Einzelaufgaben in zeitlicher Abfolge an den jeweiligen Zuständigkeitsbereich delegiert werden. Die Marketing-Fachleute eruieren die Konsumentenwünsche und -wahrnehmungen, um ein Produktkonzept zu entwickeln; die Forschungs- und Entwicklungsingenieur/-innen wählen die adäquate Technologie aus; die Produktionsingenieur/-innen entwerfen den Fertigungsprozess und andere Spezialisten tragen das Projekt als Staffelstab weiter von Phase zu Phase des Innovationsprozesses.

Anders dagegen beim Rugby-Ansatz: Der Innovationsprozess organisiert sich selbst durch die stetigen Interaktionen innerhalb eines Teams, das aus handverlesenen Mitgliedern unterschiedlicher technischer, sozialwissenschaftlicher und gestalterischer Disziplinen besteht. Ein solches Team arbeitet vom Anfang bis zum Ende des Entwicklungsprozesses zusammen. Der Fortgang

des Projekts folgt dem Rhythmus, der sich aus dem Zusammenspiel der Teammitglieder ergibt, und weniger dem Takt von zuvor festgelegten Arbeitsschritten. So beginnen zum Beispiel Ingenieure/-innen mit der Entwicklung des Produkts, bevor die Tests zur Tauglichkeit der zu verwendenden Produkttechnologie abgeschlossen sind. Ferner ist das Team gezwungen, eine Entscheidung aufgrund einer später eintreffenden Information zu revidieren und diese in einer Wiederholungsschleife in die Produktentwicklung einfließen zu lassen. Für Takeuchi und Nonaka ist dies die Ursache für einen schnelleren und vor allem flexibleren Produktentwicklungsprozess.

Der Rugby-Ansatz ist mit einem charakteristischen Führungsstil verbunden, der sich dadurch auszeichnet, größtmögliche Autonomie mit ehrgeizigen Vorgaben zu koppeln. Das neukonstituierte Team beginnt an einem Zustand Null. Nach einiger Zeit soll das Team wie ein Start-Up-Unternehmen innerhalb der Firma funktionieren, das eigene Initiativen ergreift, Risiken auf sich nimmt und seine eigenen Ziele setzt. Die damit verbundene Autonomie bedeutet, dass das Topmanagement auf Kontrollen und Interventionen weitgehend verzichtet. Es stellt lediglich die finanziellen Mittel für das Team bereit und handelt insofern wie ein *venture capitalist* im eigenen Haus. »We open up our purse but keep our mouth closed«, zitieren die Verfasser einen der befragten Vorgesetzten (ebd.: 139). Im Vergleich zum linearen Modell überlappen sich in hohem Maße die Phasen des Innovationsprozesses. Obgleich die Teilnehmer/-innen unterschiedliche Zeithorizonte benötigen – Forschung und Entwicklung beanspruchen einen weiter gefassten als das Marketing –, teilen sich sämtliche Teammitglieder ihr Wissen (*knowledge sharing*). So bekommt die/der Marketing-Experte/-in einen Eindruck von den ingenieurwissenschaftlichen Forschungsproblemen und der/die Ingenieur/-in lernt die Marktverhältnisse kennen. Durch diesen Effekt des gemeinsamen Lernens wird aus dem Team, das als eine organisierte soziale Einheit zu Anfang seine Arbeit aufnahm, eine solche, die auf Basis eines gemeinsam geteilten Wissens agieren kann. Dies erweist sich in verschiedener Hinsicht als ein Vorteil. Wissen, das im Phasenmodell separiert war, steht nun dem Team insgesamt zur Verfügung. Die Marktforschung kann

zum Beispiel durch die Informationen über die Möglichkeiten des Produktes und der zugrunde liegenden Technologie zielgerichteter ausgelegt werden, und die Ingenieure/-innen können neue Erkenntnisse über Änderungen des Konsumentenverhaltens in ihre Entwicklungsarbeit integrieren. Dadurch kann das Problem des Phasenüberganges, welches das Phasenmodell kennzeichnet, vermieden werden: An jedem Ende einer Phase eines linear organisierten Prozesses, wenn eine Arbeitsgruppe der nächsten das Projekt weiterreicht, blieb relevantes Wissen auf der Strecke. Hinzu kommt, dass solche Übergänge zum Flaschenhals werden, wenn eine Arbeitsgruppe das Zeitlimit, das für diese Phase angesetzt ist, nicht einhalten kann. Der Innovationsprozess insgesamt gerät ins Stocken. Durch die Überlappung können die Teilnehmer/-innen eines nach dem Rugby-Ansatz organisierten Teams flexibler mit Verzögerungen umgehen. Das gesamte Projekt kommt nicht zum Halt, weil die Teilnehmer/-innen der nachfolgenden Entwicklungsschritte auf der Basis der bereits vorliegenden Informationen ihre Arbeit leisten können.

Nicht zu verschweigen sind die Nachteile und die Grenzen des Rugby-Modells. Von den Teilnehmer/-innen verlangt es eine intensive Kommunikation untereinander. Eine solche Kommunikationsdynamik ist nicht nur schwer zu kontrollieren, sondern beansprucht wesentlich mehr Zeit von den Teammitgliedern als bei Innovationsprojekten, die in Phasen strukturiert sind. Außerdem ist das Rugby-Modell nicht generell für jedes Unternehmen und für jedes Innovationsziel geeignet. Insbesondere die kommunikative Abstimmung, die in einem überschaubaren Team möglich ist, stößt bei einer gewissen Anzahl von Teilnehmer/-innen an ihre Grenzen. Um den Kommunikationsaufwand zu dezimieren, ist es insbesondere bei Mammutprojekten mit einem extrem hohen Entwicklungsaufwand und einem großen Mitarbeiter/-innenstab erforderlich, die Phasensegmente des linearen Modells zu nutzen, um sinnvolle Kommunikationseinheiten zu schaffen.

3.3 Feuerwerk-Modell: Innovationsbündel und Rückschläge

Verglichen mit dem Ketten- und Rugby-Modell steht das Feuerwerk-Modell auf dem solidesten wissenschaftlichen Fundament. In Longitudinalstudien von 14 industriellen und staatlichen Projekten (darunter Projekte der »NASA« und des Unternehmens »3M«) untersuchten mehrere Forschungsgruppen des Minnesota Innovation Research Teams unter der Leitung von Andrew Van de Ven länger als ein Jahrzehnt Innovationsverläufe (Van de Ven et al. 1999). Sowohl quantitative als auch qualitative Methoden kamen dabei zum Einsatz. Die Grundthese der Forschungsgruppe um Van de Ven besagt, dass ein Innovationsprozess einer non-linearen Dynamik unterliegt. Er ist daher weder stabil und vorhersehbar, noch zufallsabhängig (ebd.: 5). Unvorhersehbar ist eine Innovation nicht deswegen, weil sie dem Zufall ausgeliefert, sondern weil sie einem komplexen Wechselspiel der am Innovationsprozess beteiligten Akteure unterworfen ist. Der Verlauf dieses Wechselspiels stellt sowohl für die Beteiligten als auch für die sozialwissenschaftlichen Beobachter ein Experiment mit offenem Ausgang dar. Auf die Metapher des Feuerwerks greifen die Autoren zurück, um auf ein essenzielles Ergebnis ihrer Forschungen hinzuweisen. Ihrer Ansicht nach verzweigt sich der Innovationsprozess unmittelbar nach der Planungsphase in viele Seitenentwicklungen. Die innovative Ausgangsidee – um die Metapher des Feuerwerks aufzugreifen – »explodiert«, um in teils auseinander driftenden, teils parallelen Bahnen fortzubestehen. Ich komme später darauf zurück. Nebenbei warnen die Autoren davor, das Handeln der an der Innovation beteiligten Akteure ausschließlich in der Annahme zu beobachten, es sei kreativ und zielkonform. Um einen positiven Bias in der sozialwissenschaftlichen Beobachtung zu vermeiden, ist es sinnvoll, deren Handeln ebenso nach Gesichtspunkten zu untersuchen, inwieweit es Innovationsprojekte behindert oder gar ganz unterbindet.

Bevor ich dieses Prozessmodell erläutern werde, ist es unerlässlich, auf jene Befunde einzugehen, welche hierzu die Grundlage bilden. Bei ihren Untersuchungen konzentriert sich die Forschungsgruppe von Van de Ven auf bestimmte Variablen. Darun-

ter fallen Ideen, Akteure, Transaktionen und Kontexte. Zu Beginn der Studie war die Ansicht weit verbreitet, dass eine singuläre Idee (zum Beispiel eine Erfindung) durch den Entwicklungsprozess entfaltet wird, während die Identität der Ausgangsidee unverändert bleibt, so als ob in dieser *in nuce* die spätere Innovation enthalten sei. Dagegen wenden Van de Ven et al. ein, dass eine Innovation selten auf einem einzigen Kernkonzept beruht. Ihren empirischen Ergebnissen zufolge bringen Organisationen zeitgleich ein ganzes Bündel von innovativen Ideen auf zum Teil divergierenden Entwicklungsbahnen voran (ebd.: 10). Während Schumpeter und seine Nachfolger den heroischen Unternehmer (*entrepreneur*) oder das Einzelunternehmen als den entscheidenden Akteur des Wandels identifizieren, betonen Van de Ven und sein Team, dass die meisten Innovationen zu komplex sind, als dass sie von einer Person allein in die Wege geleitet werden können. Eine Innovation emergiert vielmehr in der Interaktion zwischen Führungskräften und Mitarbeitenden verschiedener Spezialgebiete, Erfahrungen und Fertigkeiten sowie darüber hinaus in den Interaktionen zwischen verschiedenen Organisationen (vgl. Kap. IV). Mit dieser Fokussierung auf die interaktiven Aspekte des Innovationsverlaufs rücken insbesondere die Beziehungen zwischen den organisierten sowie personifizierten Akteuren in den Blick. Solche multilateralen Beziehungen, die Van de Ven et al. untersuchen, entwickeln sich nicht nach dem sequenziellen Schema ›Verhandlung – Einigung – Ausführung‹. Stattdessen werden unterschiedliche Phasen von hoher und niedriger Interaktivität beobachtet, die nach kurzer Zeit den Rahmen einer dyadischen Beziehung zwischen zwei Organisationen überspringen und sich zu einem Netzwerk von Beziehungen auswachsen.

Konsequenterweise erweitern die Autoren die für Manager/-innen primäre Sicht auf den betriebswirtschaftlichen Rahmen, indem sie die gesellschaftliche Umwelt des Unternehmens beleuchten. Es handelt sich hierbei um jene gesellschaftlich-technologische Infrastruktur, bestehend aus technischen Normen, wissenschaftlichen Expertisen, zugänglichen Testlaboratorien, Kapitalgebern, Fachkräften ausbildende Hochschulen und nicht zuletzt informierten Konsumenten, die mit Innovationen etwas anzufangen wissen. Es gilt dabei anzuerkennen, dass eine solche In-

Infrastruktur zu einem großen Teil als ein öffentliches Gut (*public good*) entwickelt wird, welches Unternehmen in ein privates Gut (*private good*) überführen können. Diese gesellschaftlich-technologische Infrastruktur stellt nicht eine vom Unternehmen losgelöste Umwelt dar, sondern kann durchaus mitgestaltet werden. Voraussetzung hierfür ist jedoch, dass ein Unternehmen ein Management (*corporate governance*) praktiziert, das sowohl spezialisiert darauf ist, diese Infrastruktur zu nutzen, als auch dazu fähig, diese mitzugestalten.

Den Innovationsprozess selbst charakterisieren Van de Ven et al. mit dem romantischen Bild der Reise. Damit soll zum Ausdruck gebracht werden, dass jedwede Innovation ihrem eigenen Weg folgt und streckenweise auf ungeplanten und unplanbaren Routen am Ziel ankommt (oder auf der Strecke bleibt). Trotz der partiellen Unvorhersehbarkeit des Verlaufs behaupten die Autoren, in den von ihnen untersuchten Fällen ein Muster zu erkennen und eine generalisierende Definition wagen zu können: »The innovation journey is a nonlinear cycle of divergent and convergent activities that may repeat over time and at different organisational levels if resources are obtained to renew the cycle« (ebd.: 16). Auf der Basis ihrer Studienergebnisse entwerfen Van den Ven und seine Mitarbeiter eine Art Landkarte (*roadmap*) des Feuerwerk-Modells, um die typischen Verlaufsformen eines Innovationsprozesses festzuhalten. Vorweg ist anzumerken, dass es sich bei den von ihnen untersuchten Fällen allesamt um »generische« Innovationen handelt. Sie zeichnen sich durch absichtsvolle Anstrengungen, eine Idee zu verwirklichen aus, beinhalten organisatorische, technologische und ökonomische Formen der Unsicherheit, bedürfen einer kollektiven Anstrengung über einen längeren Zeitraum und erfordern mehr Ressourcen, als die Innovatoren selbst in den Händen halten (ebd.: 22).

Dem Modell zufolge beginnt der Innovationsprozess nicht mit der Inspiration eines einzelnen Akteurs, sondern verläuft eher schleppend in einer über mehrere Jahre ausgedehnten Periode der »Schwangerschaft«, bis mehr und mehr ähnliche Ideen sichtbar werden. Diese Etappe verläuft evolutionär und weitgehend ohne dezidierte Planung. Nichtsdestotrotz ist eine Schlüsselfigur auszumachen, in der die scheinbar losen Fäden zusammenlau-

fen. In günstigen Momenten bieten diese Schlüsselfiguren ihr eigenes Unternehmen oder ihre Abteilung als eine Art ›Gefährter‹ an, um die Ideen weiterzuentwickeln und Krisen zu überwinden. Erst Schocks, wie eine Budgetkrise, ein neues Management, ein Markteinbruch etc., lösen konzentrierte Bemühungen der Akteure aus. Die Wirkung des Schocks besteht darin, diejenige Aufmerksamkeit, die sich im Stadium der Routine auf das Bewährte richtet, auf das Neue zu lenken. Die Krise, die Suche nach Lösungen motivierend, erhöht die Wahrscheinlichkeit, sich mit den Konzepten der Schlüsselfiguren eingehender auseinander zu setzen. Der offizielle Innovationsprozess beginnt mit der Entwicklung von Plänen, um Vorgesetzte oder potenzielle Kapitalgeber/-innen für das Projekt zu gewinnen.

Bemerkenswert an den Studien von Van de Ven et al. (1999) ist der Befund, dass in sämtlichen untersuchten Fällen die finanziellen, personellen und technologischen Ressourcen für die Innovation von außen in die betreffenden Unternehmensabteilung gelangen. Bei der Planung geht es in erster Linie darum, die Informationsasymmetrie zwischen den Innovatoren und den Kapitalgebern zu überwinden (›principal-agent«-Problem). Plänen sind daher eher eine argumentativ-rhetorische Funktion zugeordnet, Zweifler umzustimmen, als dass sie als realistisches Szenario für ein Innovationsprojekt gelten können. Daher unterschätzen sämtliche Pläne den Kosten- und den Zeitaufwand und übergehen geflissentlich mögliche Unsicherheiten. Für diese Realitätsverzerrung gibt es zwei Gründe: *erstens* die bereits erwähnte rhetorische Funktion. Es gilt, potenzielle Unterstützer/-innen zu überzeugen – darauf spekulierend, dass, wenn jemand Kapital in das Projekt investiert hat, er dann auch bereit ist, die Folgekosten zu tragen, um seinen bereits geleisteten Einsatz nicht zu verspielen. *Zweitens* handelt es sich bei Plänen nicht selten um Selbsttäuschungen von Menschen, die sich aufgrund ihrer Identifikation mit dem Projekt und der emotionalen Involviertheit den nicht unwahrscheinlichen Misserfolg nicht ausmalen wollen und können, weil eine solche Vorstellung zudem entmutigen würde (ebd.: 33).

Während das herkömmliche lineare Modell davon ausgeht, eine Projektidee bliebe bei der Weiterentwicklung intakt und werde

lediglich den Umständen entsprechend angepasst, geht das Feuerwerk-Modell davon aus, dass die Anfangsidee gleich nach der Planung in multiple Entwicklungen zerfällt. Da das ›Feuerwerk‹ den Beteiligten den Eindruck vermittelt, die Ausgangsidee sei fruchtbar, kommt es bei den Beteiligten zu einer Anfangseuphorie (*honeymoon*), die alsbald jedoch in einem mehr oder weniger kontrollierten Chaos mündet. Sobald die Entwicklungsarbeiten beginnen, ist zu beobachten, wie sich die Ausgangsidee in zahlreichen Versionen fortspinnt, sei es auf divergenten, sei es auf parallelen oder konvergenten Wegen. Diese Aufspaltung in verzweigte Aktivitäten gehört zu den bisher am wenigsten verstandenen Merkmalen des Innovationsprozesses. Van de Ven und seine Mitautoren nennen mögliche Gründe für dieses ›Feuerwerk‹ (ebd.: 35). Zu den wichtigsten gehört, dass sich eine Innovation *erstens* als eine vieldeutige und unsichere Angelegenheit erweist und es oftmals unmöglich ist, anfangs zu wissen, welche Abzweigung zum Ziel führt. Daher ist es sinnvoll, alternative Entwicklungen zuzulassen und zu explorieren. *Zweitens* können die Verzweigungen des Feuerwerks dazu dienen, das Risiko des Innovationsprozesses auf verschiedene Einzelentwicklungen zu verteilen.

Häufig ereilen Rückschläge den Prozess. Mal gehen Pläne nicht auf, ein anderes Mal verändern Umweltereignisse (extreme Schwankungen von Rohstoffpreisen, Ausscheiden von Partnern oder tragenden Mitarbeiter/-innen etc.) die Voraussetzungen für den Erfolg. Wenn Zeit- und Etatpläne merklich auseinander klaffen, hat die Gnadenfrist für das Projekt begonnen. Rückschläge wirken sich dann ungünstig aus, wenn sie, was häufig der Fall ist, zu Domino-Effekten führen. Parallele Innovationsaktivitäten werden solange als unproblematisch erachtet, wie sie voranschreiten. Sobald jedoch ein Teilprojektziel verfehlt wird, wird den Beteiligten bewusst, wie stark die Abhängigkeiten voneinander sind. Spätestens nach der Krise werden die Erfolgskriterien neu ausgehandelt. Dabei ist der Konflikt zwischen den Innovatoren und den Ressourcengebern vorprogrammiert. Während Erstere den Rückschlag als Zeichen dafür interpretieren, dass der eingeschlagene Weg nicht konsequent genug verfolgt wurde, ist es für Letztere ein Anlass, die Investitionen zu überdenken.

Nach vielen Etappen, von denen ich nur die wichtigsten angesprochen habe, erreicht die Reise ihr Ziel: Die Innovation wird entweder implementiert oder bricht ab. Bei der Implementation geht es darum, in inkrementellen Anstrengungen das Neue mit dem Alten zu verbinden. Wenn Investoren und Vorgesetzte die Entscheidung zum Abbruch des Innovationsprozesses treffen, dann attribuieren sie die Verantwortung für das Scheitern zumeist an die Person des Innovators (und nicht an eventuell ungünstige Randbedingungen). Auf diese Weise wird der Innovator stigmatisiert; seine Chancen, im selben Unternehmen erneut einen Innovationsprozess anzustoßen, sinken.

4. ZWISCHENRESÜMEE

Kurz bilanzierend will ich die linearen mit den non-linearen Modellen auf einer generellen Ebene kontrastieren. Der Unterschied besteht *erstens* darin, dass non-lineare Modelle ein System unterschiedlicher und rekursiv verknüpfter Ursache-Wirkungs-Ketten repräsentieren und sich nicht auf eine einzige lineare beschränken. *Zweitens* betont die Mehrzahl der non-linearen Modelle die Diskontinuität in der technischen Entwicklung. Sie beschreiben diese als ein Wechselspiel von Perioden steten Fortschreitens und innovatorischen Durchbrüchen. *Drittens* zeigen die non-linearen Modelle, dass sich zwar unterschiedliche Phasen des Innovationsprozesses markieren, jedoch nur unzureichend abgrenzen lassen. Vereinfacht man die Sicht und gliedert einen Innovationsprozess in säuberliche Zeitsegmente, so unterschlägt man damit all die Rückkopplungen und Rückschläge, die entscheidend für die Entwicklung einer Innovation sein können.

Viele gesellschaftlich bedeutsamen technischen Innovationen – man denke zum Beispiel an das Telefon oder an den Computer – nahmen erst, nachdem sie mehrere Umwege hinter sich gebracht hatten, ihre heute akzeptierte Form an; und selbst diese Form ist nur vorübergehend fixiert, da Re-Inventionen eingeführte Techniken ständig umformen können. Innovationen zeichnen sich oft durch ihre fluide Identität aus. Das Objekt am Ende eines Innovationsprozesses ist mit dem, was anfänglich intendiert, ge-

plant und entworfen wurde, selten identisch. Eine Ursache für diese Wandlungsfähigkeit liegt darin, dass die Auseinandersetzung mit neuen Ideen bei den beteiligten Akteuren einen Lernprozess anstößt, der dazu führt, Produkte, Designs, Vorstellungen etc. zu verändern und den konkreten Bedürfnissen und Interessen anzupassen. Eine Kernaussage des letzten Kapitels lautete: Es gibt keinen ›Direktzubringer‹ von der Erfinderwerkstatt hin zu den Orten der Nutzung. Insbesondere das Feuerwerk-Modell betont die Dynamik des Innovationsprozesses, der oftmals von multiplen Ideen und Möglichkeiten geprägt ist. Selbst dann, wenn sich Umrisse einer Innovation verfestigen, nimmt die Innovation selten ihren vorgezeichneten Gang. Viel eher sind kritische Passagen, Rückschläge, Sprünge, manchmal sogar jahrzehntelange ›Scheintodphasen‹ zu beobachten, im Laufe derer die Akteure jegliche Hoffnung auf einen Erfolg fahren gelassen haben. Innovationsprozesse folgen demnach einer Logik, die sich weder als eindeutig vorhersehbar und berechenbar erweist, noch als vollkommen zufällig.

Nur im eng gesteckten Rahmen von inkrementellen Innovationen liefern Phasenmodelle dem Management ›die eine‹ hinreichend stabile Struktur, Innovationsprozesse planen und steuern zu können. Ansonsten ist aus meiner Argumentation abzuleiten, wie schwierig es ist, generelle und verbindliche Ursache-Wirkungs-Modelle für Innovationen anzubieten, deren Verlauf zu prognostizieren, diese betriebswirtschaftlich durchzukalkulieren und politisch zu steuern. Es ist einzuräumen, dass die vorgestellten Prozesstheorien nur in Ausnahmefällen dem Anspruch gerecht werden, Erklärungsmodelle zu sein, wie wir es von den exakten Naturwissenschaften erwarten. Doch das Verfehlen dieses Ideals sollte uns nicht davon abhalten, um die Entdeckung von Regelmäßigkeiten, Muster und Bedingungsfaktoren zu ringen, die uns in wissenschaftlicher, politischer und wirtschaftlicher Hinsicht einen rationaleren Umgang mit Innovationen gestatten, als dies ohne Rückgriff auf diese Einsichten möglich ist.