

Mikrofundierung von Forschungs- und Entwicklungsausgaben der Firma bei unterschiedlichen Marktstrukturen

Ein Überblick¹

Von Jörg Flemmig

Neue Technologien fallen weder als Manna vom Himmel noch werden sie hinreichend aus dem größeren Erfahrungsschatz einer erweiterten Produktion erklärt: Sie müssen produziert werden. Die Höhe der Ausgaben für Forschung und Entwicklung ($F + E$) wird endogen aus den Kalkülen der Firma bestimmt, die sich gewinnmaximierend verhält, indem sie unter Berücksichtigung der mit einer Innovation verbundenen Einnahmen und Kosten und der Marktstruktur, in der sie agiert, den gewinnmaximalen Entwicklungszeitpunkt der technischen Neuerung wählt.

1.0. $F + E$ Kalküle einer Firma bei nahezu vollkommener Konkurrenz

In diesem Abschnitt wird eine Marktstruktur unterstellt, die sich von einem vollkommenen Wettbewerb nicht wesentlich unterscheidet; die einzelne Firma ist lediglich unsicher über den Innovationszeitpunkt der Konkurrenz. Unsicherheit wird konzeptionell im Rahmen subjektiver Wahrscheinlichkeitsannahmen der Firma bezüglich des Verhaltens der Konkurrenz behandelt. Würde der Markt den Gesetzmäßigkeiten der vollkommenen Konkurrenz folgen, ließen sich Innovationen systematisch kaum erklären, da bei Sicherheit und unendlich schneller Marktanpassung jede innovative Firma sofort mit Imitation durch die Konkurrenz rechnen müßte. Es sei $P(t)$ die Wahrscheinlichkeitsannahme einer Firma zum Zeitpunkt $t = 0$, daß bis zum Zeitpunkt $t > 0$ die Konkurrenz mit einem neuen Produkt oder Produktionsverfahren auf dem Markt ist. In der Literatur wird dabei kaum zwischen einer Verbesserung des Produktionsverfahrens und einer Produktinnovation differenziert. Im ersteren Fall „verschiebt“ sich die Produktionsfunktion nach oben, die Kosten werden gesenkt. Ein wirklich neues Produkt, z.B. ein Lasergerät, erweitert den Güterraum der Ökonomie um eine Dimension und impliziert eine völlig neue Produktionsfunktion.

Die Klassifikation einer Innovation als Produkt- oder Verfahrensinnovation wird in der Praxis oft eine Frage der Betrachtungsweise sein. Ein neuer

¹ Für zahlreiche Hinweise habe ich den Professoren Dr. W. Harbrecht, Universität Passau, Dr. W. Vogt, Universität Regensburg, und zwei anonymen Gutachtern zu danken.

Computer ist sicher für dessen Hersteller eine Produktinnovation, der Käufer des Computers wird jedoch mehr an den kostensenkenden Aspekten des durch die Anwendung des Computers möglichen neuen Produktionsverfahrens interessiert sein.

In der Literatur wird Unsicherheit häufig mit Hilfe der Hazardrate $h(t)$ in die Innovationsmodelle integriert, die die bedingte Wahrscheinlichkeit dafür angibt, daß die Konkurrenz im nächsten Augenblick auf den Markt kommt, wenn ihr dies bis zu diesem Zeitpunkt $[t, t + dt]$ nicht gelungen ist².

$$(1) \quad h(t) = \frac{P'(t)}{1 - P(t)}.$$

$P(t)$ wird üblicherweise in der Literatur durch folgende Exponentialverteilung konkretisiert:

$$(2) \quad P(t) = 1 - e^{-ht} \text{ für } t \geq 0.$$

Berücksichtigt man die entsprechende Dichtefunktion in (1), erhält man eine konstante Hazardrate h .

Da der Erwartungswert der Verteilung $P(t) = 1 - e^{-ht}$ im Zeitpunkt $t = 0$ gleich $1/h$ ist, ist für die Firma der Zeitpunkt bekannt, an dem sie mit der Einführung einer Innovation durch die Konkurrenz rechnen muß.

1.1. Die Einnahmenstruktur einer innovativen Firma

Die Firma bezieht bis zur Einführung einer Innovation aus dem Verkauf eines Produktes einen Nettoeinkommensstrom V_0 pro Zeiteinheit, der über einen bestimmten Zeitraum mit der konstanten Rate g zu- oder abnehmen kann, je nachdem, ob der Markt für dieses Gut mit der Rate g wächst oder schrumpft.

Der Nettoeinkommensstrom der Firma ändert sich, wenn sie zum Zeitpunkt $T \geq 0$, $t = 0$ sei die Gegenwart, durch ein technisches Verfahren oder ein neues Produkt die Kosten senkt bzw. das alte Produkt obsolet wird. Es handelt sich bei der Innovation nicht um ein völlig neues Gut. Sie erhält bei einer Innovation vor der Konkurrenz ab dem Zeitpunkt T Einnahmen in Höhe von V_2 pro Zeiteinheit.

Für die Firma, die zum Zeitpunkt T eine Innovation auf den Markt bringt, wird sich der Nettoeinkommensstrom ändern, wenn andere Firmen die Erfindung imitieren oder selbst mit einem neuen Produkt auf den Markt kommen. Die neuen Einnahmen pro Zeiteinheit belaufen sich dann auf V_3 .

² Vgl. *Sivazlion / Stanfel* (1975), *Loury* (1979), *Lee / Wilde* (1980), *Kamien / Schwarz* (1982).

Schließlich bleibt für die Firma noch die Möglichkeit, daß sie mit der Innovation erst nach den Konkurrenten auf den Markt kommt, also selbst Imitator ist. Dies bedeutet für die Firma zunächst eine Reduktion ihrer Einnahmen pro Zeiteinheit von V_0 auf V_1 bis zum Zeitpunkt ihrer Imitation der Innovation, ab dem sie mit Einnahmen in Höhe von V_4 pro Zeiteinheit rechnen kann. Weitere Innovationszyklen werden nicht berücksichtigt. Plant die Firma also eine Innovation zum Zeitpunkt T und rechnet sie bei gegebenem Zinssatz r mit den oben beschriebenen Einkommensströmen, Marktbewegungen und Wahrscheinlichkeiten, ergibt sich folgendes erwartetes Einkommen $v(t)$:

$$(3) \quad v(t) = \int_0^T e^{-(r-g)t} \left[V_0 (1 - P(t)) + V_1 P(t) \right] dt + \int_T^\infty e^{-(r-g)t} \left[V_2 (1 - P(t)) + V_3 (P(t) - P(T)) \right] dt + \int_T^\infty e^{-(r-g)t} V_4 P(T) dt.$$

Die möglichen Einnahmengrößen werden mit ihrer Eintrittswahrscheinlichkeit gewichtet und unter Berücksichtigung der Entwicklung der Marktgröße diskontiert³.

1.2. Die Kostenstruktur von Forschung und Entwicklung

Die Kosten C eines Forschungsprojektes hängen ab von der verwendeten Technologie, der erwünschten Qualität des neuen Produktes/Produktionsverfahrens und der Zeitdauer, innerhalb der man das Projekt erfolgreich abschließen will.

Durch Variation des Ausgabenstromes für $F + E$ pro Zeiteinheit können Kosten gespart werden, wenn die Ausgaben später erfolgen, da dann der Gegenwartswert der Forschungsaufwendungen geringer ist, jedoch verschiebt sich damit auch der Innovationszeitpunkt T ; die gesamten Kosten werden steigen, wenn aufgrund strategischer Überlegungen die Innovation zu einem früheren Zeitpunkt erfolgen soll. Die Entwicklungszeit kann z.B. dadurch verkürzt werden, daß mehr Leute an einem Projekt arbeiten, mehrere Projekte parallel laufen, von denen sich einige als unnütz erweisen werden, oder einfach Zwischenergebnisse nicht hinreichend überprüft werden, was die Fehlerquote erhöhen wird.

³ Tritt eine Firma neu in den Markt ein, kann sie bis zu einer Innovation nicht mit laufenden bzw. reduzierten Einnahmen rechnen: $V_0 = V_1 = 0$.

Unterstellt man zunächst Sicherheit bezüglich der Kosten, die bei einer geplanten Innovationszeit T entstehen, kann man davon ausgehen, daß die Kosten C mit zunehmender Rate steigen, wenn man die Entwicklungszeit verkürzt. Dies impliziert:

$$(4) \quad c(T) > 0; \quad c'(T) < 0; \quad c''(T) > 0.$$

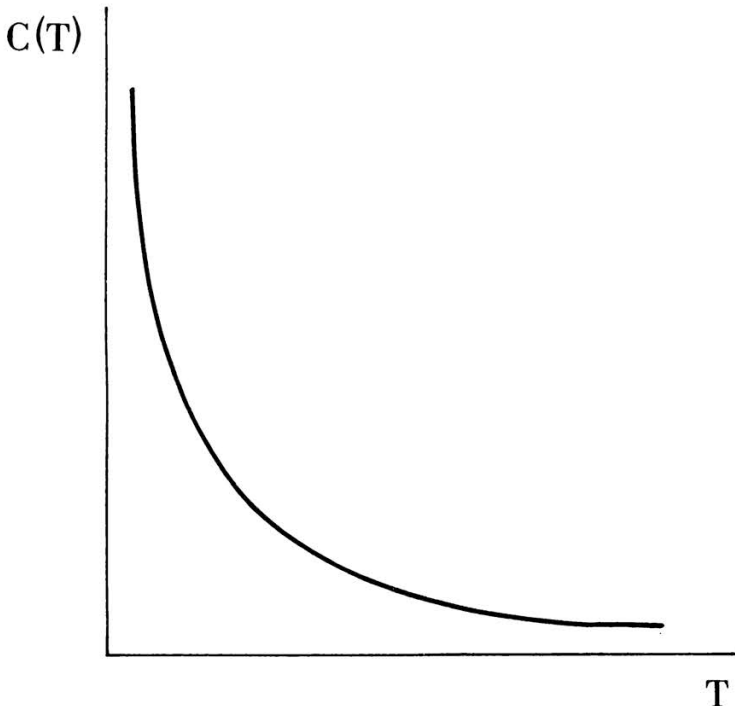


Schaubild 1

1.3. Der gewinnoptimale Entwicklungszeitpunkt der Firma

Aufgrund der angenommenen Einnahmen- und $F + E$ Kostenstruktur kann nun der gewinnoptimale Entwicklungszeitpunkt ermittelt werden. In ihm wird die Differenz aus erwarteten Einnahmen und erwarteten Kosten optimiert.

Zunächst wird angenommen, daß im Zeitpunkt $T = 0$ die geplanten Ausgaben für $F + E$ durch Verträge fixiert werden, so daß die Ausgabenstruktur für $F + E$ nicht modifiziert werden kann, wenn ein Konkurrent vor dem geplanten Entwicklungszeitpunkt T mit der Innovation auf den Markt kommt.

Der Entwicklungszeitpunkt T ist optimal, wenn die Grenzeinnahmen im Zeitpunkt $[T + dt]$ gleich den Grenzkosten sind.

Durch eine kleine Verzögerung entstehen marginale Nettoeinnahmen durch die Reduzierung von Kosten; die marginalen Verluste hängen davon ab, ob ein Konkurrent bis zum Zeitpunkt T auf den Markt gekommen ist. Ist dies der Fall, impliziert ein Verzicht auf die Innovation weiterhin reduzierte Einnahmen (V_1) im Zeitpunkt $[T + dt]$, während die Firma bei einer Innovation mit Einnahmen aus der Imitation (V_4) in diesem Zeitpunkt rechnen kann; ist die Konkurrenz im Zeitpunkt $[T + dt]$ noch nicht auf dem Markt, ergibt sich der marginale Verlust aus der Differenz zwischen den Einnahmen bei Innovation und dem Einkommenstrom aus dem laufenden Geschäft in diesem Zeitpunkt ($V_2 - V_0$).

Kommt die Konkurrenz gerade im Zeitpunkt $[T + dt]$ auf den Markt, verliert man den Einnahmenstrom eines Innovators, der später imitiert wird und wird selbst Imitator.

Intuitiv einsichtig ist der Einfluß der Parameter V_0 , V_1 , V_2 , V_3 , auf den gewinnoptimalen Entwicklungszeitpunkt T :

- Nehmen die laufenden Einnahmen (V_0) zu, wird die Innovation erst später erfolgen.
- Erhöht sich der Einnahmenstrom, der der Firma nach Erscheinen der Konkurrenz verbleibt, wenn sie selbst die neue Technologie noch nicht entwickelt hat (V_1), wird sie einen späteren Entwicklungszeitpunkt wählen.
- Steigt der Einnahmenstrom aus der Innovation (V_2), erfolgt die Innovation früher.
- Eine Erhöhung der Einnahmen, wenn man selbst imitiert wird (V_3), führt zu einer Vorverlegung des Innovationszeitpunktes.
- Erhöht sich die Imitationsgeschwindigkeit der Konkurrenz, reduzieren sich die marginalen erwarteten Einnahmen, was durch eine Reduzierung der erwarteten Entwicklungskosten ausgeglichen werden muß: Die Entwicklungszeit wird verlängert.

Sind die Entwicklungskosten nicht vertraglich fixiert, kann die Firma ihre F + E Ausgaben sofort stoppen, wenn die Konkurrenz die Innovation zuerst auf den Markt bringt, oder aber die Forschung auf neuem Niveau weiterführen, um den Einnahmenstrom aus der Imitation zu erhalten. Die Ausgaben sind also nicht durch die Wahl des gewinnoptimalen Entwicklungszeitpunktes determiniert; sie sind insoweit stochastisch, als zwar bis zum geplanten Entwicklungszeitpunkt T ein bestimmter Ausgabenstrom $x(t)$ pro Zeiteinheit erfolgt, dies aber nur, wenn die Konkurrenz nicht schneller auf den Markt kommt, also mit der Wahrscheinlichkeit $(1 - P(t))$.

Der Gegenwartswert K_0 der erwarteten Kosten einer Innovation zum Zeitpunkt T beträgt

$$(5) \quad K_0 = \int_0^T e^{-rt} x(t) (1 - P(t)) dt.$$

Der Gegenwartswert der erwarteten Einnahmen wird affiziert vom Gegenwartswert (W) eines Patentes im Zeitpunkt T , den die Firma mit der Wahrscheinlichkeit $(1 - P(t))$ erhält, den erwarteten Einnahmen aus den laufenden Geschäften vor einer eigenen oder fremden Innovation ($V_0 (1 - P(t))$) und dem Wert des maximalen Nettorückflusses S , der für die Firma noch möglich ist, wenn die Konkurrenz vor dem von der Firma geplanten Entwicklungszeitpunkt T auf den Markt kommt. Der Gegenwartswert der erwarteten Nettoeinnahmen V beträgt bei einer Wachstumsrate des Marktes in Höhe von g

$$(6) \quad V = \int_0^T e^{-rt} \left[e^{gt} V_0 (1 - P(t)) + SP'(t) \right] dt + e^{-(r-g)T} W (1 - P(T)).$$

Will die Firma den erwarteten Gewinn über die Wahl des Entwicklungszeitpunktes T maximieren, muß sie insbesondere den Wert ihres maximalen Nettorückflusses kennen, mit dem sie ab dem Entwicklungszeitpunkt der Konkurrenz k_0 , $k_0 < T$, noch rechnen kann: Die Planung im Zeitpunkt $t = 0$ setzt die Analyse der Situation im Zeitpunkt $k_0 < T$ voraus.

Die Firma plant einen Innovationszeitpunkt T und damit verbundene Ausgaben, an denen sie bis zum Innovationszeitpunkt festhält, solange kein Konkurrent mit der Entwicklung schneller ist; bringt jedoch die Konkurrenz die Innovation eher auf den Markt, wird die Firma ihren Ausgabenplan und damit den geplanten Entwicklungszeitpunkt möglicherweise revidieren.

Es sei W_m der Gegenwartswert einer Imitation im Zeitpunkt T_m für die Firma und V_1 die reduzierten Nettoeinnahmen pro Zeiteinheit, die der Firma nach dem vorzeitigen Erscheinen der Konkurrenz noch verbleiben.

Die Firma wählt einen revidierten Ausgabenplan $x_r(t)$ und modifizierten Entwicklungszeitpunkt T_m , der die noch möglichen Einnahmen V_r maximiert.

$$(7) \quad V_r = \max_{T_m, x_r(t)} \int_{k_0}^{T_m} e^{-rt} (V_1 - x_r(t)) dt + e^{-rT_m} W_m.$$

Dies jedoch unter Berücksichtigung der ab dem Zeitpunkt k_0 für die Fertigstellung der Innovation noch notwendigen finanziellen Aufwendungen, die der Firma bei Sicherheit bezüglich der Technologie ja bekannt sind.

- Die Firma beendet das Forschungsvorhaben, wählt also keinen neuen Innovationszeitpunkt T_m , wenn der Wert der Imitation gegenüber den noch notwendigen Aufwendungen relativ gering ist. Sind die reduzierten Einnahmen V_1 positiv, bleibt sie im Markt: Der Gegenwartswert ihrer noch möglichen Einnahmen im Zeitpunkt k_o beträgt dann $S = V_1/r$.

Wenn nun die Firma auf dem Markt bleibt, kann sie die Entwicklungszeit verändern: Waren es zunächst hohe erwartete Einnahmen aus der Innovation, die sie zu einer schnelleren Entwicklung motiviert haben, fällt dieser Grund bei vorzeitiger Innovation durch die Konkurrenz weg, die Firma reduziert die Entwicklungskosten durch eine Verlängerung der Entwicklungszeit; war der Anreiz für eine zügige Innovation die Angst vor hohen Einbußen, wird nach einem vorzeitigen Auftauchen der Konkurrenz die Entwicklungszeit beschleunigt, um die Dauer der Verluste gering zu halten.

Da Neuerungen nicht unbedingt bedeutende Vorteile gegenüber der laufenden Produktion bringen und eventuell auch aus Rücklagen finanziert werden können, ist der Einfluß der Budgetrestriktion im Rahmen von $F + E$ nicht so dominant; es sei denn, man ist neu auf dem Markt und kann nicht auf laufende Einnahmen zurückgreifen.

Unberücksichtigt bleibt in diesem Modell der Kredit: Hohe erwartete Einnahmen könnten als Sicherheit für Kredite dienen und würden dann sehr wohl über erhöhte Anfangskassenbestände oder höhere laufende Einnahmen während der Entwicklungsperiode die optimale Entwicklungszeit affizieren. Fatal für den Kreditnehmer ist es aber, sein Wissen um die Innovation mindestens teilweise preisgeben zu müssen, wenn er sonst keine Sicherheiten anbieten kann. Der Erfinder muß einen für ihn optimalen trade-off zwischen der Reduzierung des Wertes seines Informationsvorsprunges und der Erhöhung seiner Finanzierungsmittel suchen.

2.0. Technischer Fortschritt im symmetrischen Cournot-Nash-Gleichgewicht

Es wurde gezeigt, daß in symmetrischen Nashgleichgewichten, also bei identischen Firmen ohne die Annahme technischer Unsicherheit, es bei spieltheoretischen Kalkülen entweder keine oder höchstens eine innovative Firma gibt⁴. Bei technischer Sicherheit weiß jede Firma, welche Kosten mit einer technischen Entwicklung bis zu einem bestimmten Zeitpunkt verbunden sind; sie kennt zudem den in die Forschung investierten Betrag der Konkurrenz und geht davon aus, daß die anderen Marktteilnehmer ihren $F + E$ Aufwand nicht ändern, wenn sie selbst die Investitionsentscheidung trifft

⁴ Vgl. Dasgupta / Stiglitz (1980).

(Cournotannahme). Jede einzelne Firma würde nun eine Strategie wählen, durch die die Innovation für sie zu einem bestimmten Zeitpunkt gesichert ist, sie wäre aber durch die Konkurrenz gezwungen, den Innovationszeitpunkt so zu wählen, daß der Gegenwartswert der verwendeten $F + E$ Mittel gleich ist dem Gegenwartswert der erwarteten Einnahmen. Wenn aber jede Firma die Innovation für sich sichern kann und die Einnahmen nicht proportional mit der Anzahl der forschenden Firmen wachsen, entsteht für jede einzelne Firma bei einer Innovation ein Verlust, da jede Firma nur einen Teil der möglichen Einnahmen aus einer Innovation erhält. Da bei spieltheoretischen Ansätzen diese Kalküle erfolgen, bevor in $F + E$ investiert wird, verzichtet jede Firma auf die Innovation.

Nur wenn eine Firma diesen Mechanismus mit in ihr Kalkül einbezieht, was aber die Annahme der Symmetrie verletzen würde, kommt es zu einer Innovation, da diese cleverere Firma von einem Verzicht der Konkurrenz auf die Innovation ausgehen könnte.

Marktunsicherheit wird in diesen Modellen durch technische Unsicherheit generiert.

Dabei wird von einer stochastischen Innovation ausgegangen, bei der der Innovationserfolg bis zum Zeitpunkt t für Firma _{i} mit einer bestimmten Wahrscheinlichkeit $P(t)$ eintritt, die von einer Pauschalausgabe x_i zu Beginn der Planung ($t = 0$) abhängt⁵. Konkretisiert wird diese Wahrscheinlichkeitsverteilung in der angelsächsischen Literatur durchgängig durch folgende Exponentialfunktion:

$$(8) \quad P(t_i \leq t) = 1 - e^{-h_i(x_i)t},$$

wobei t_i den Innovationszeitpunkt, x_i die $F + E$ Ausgaben der Firma _{i} und h_i die „Hazardrate“ der Firma _{i} in Abhängigkeit von den Forschungsausgaben x_i repräsentieren. Die Hazardrate h_i gibt die bedingte Wahrscheinlichkeit für einen Innovationserfolg von Firma _{i} im Zeitintervall $[t, t + dt]$ an, wenn bis dahin noch keine Innovation erfolgt ist:

$$(9) \quad h(x_i) = \frac{P'(t, x_i)}{1 - P(t, x_i)}.$$

Die Entscheidungsvariable ist dann für jede Firma die Höhe der Forschungsausgaben, die sie investiert. Die Funktion $h(x)$ wird quasi als Produktionsfunktion des technischen Fortschritts mit zunächst zunehmenden und dann abnehmenden Skalenerträgen interpretiert. Dies wird anhand folgender Graphik verdeutlicht:

⁵ Vgl. *Loury (1979), Dasgupta / Stiglitz (1980)*.

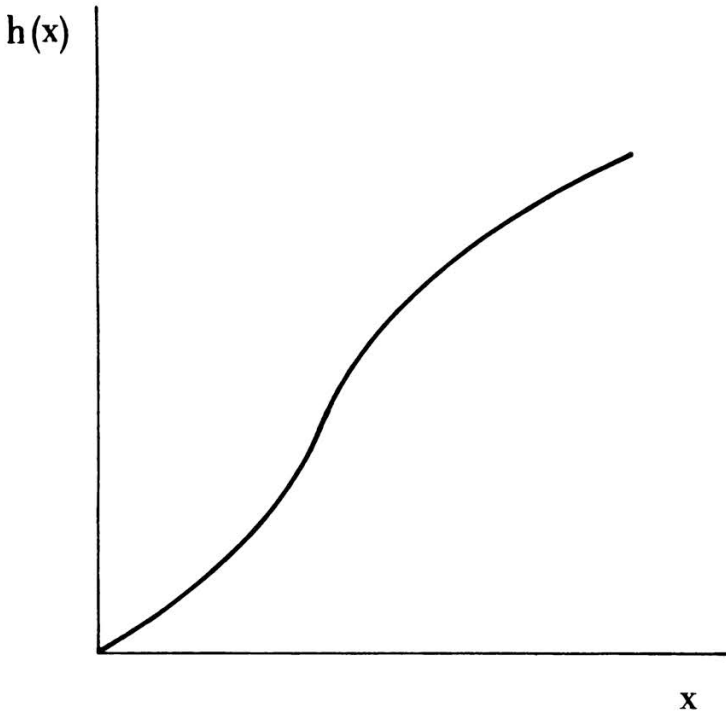


Schaubild 2

Die Einnahmen aus einer Innovation werden durch Patente abgesichert. Da ein symmetrisches Nashgleichgewicht zugrunde gelegt wurde, die n -te Firma kennt die F + E Ausgaben der $n - 1$ anderen Firmen, die unabhängig von der Ausgabenwahl der n -ten Firma sind, kann von identischen Optimierungskalkülen ausgegangen werden. Eine repräsentative Firma maximiert über die Wahl der F + E Ausgaben (x) die Differenz aus dem Gegenwartswert der erwarteten Einnahmen und sicheren Kosten, erhält also als Gewinn G

$$(10) \quad G = \max_x \int_0^{\infty} e^{-rt - nh(x)t} \pi h(x) dt - x,$$

wobei durch π die bei einer Innovation möglichen Einnahmen beschrieben werden und $e^{-nh(x)t}$ bei Verwendung von (8), die Wahrscheinlichkeit für die repräsentative Firma angibt, daß genau sie im Zeitpunkt $[t + dt]$ bei Forschungsausgaben in Höhe von x die Erfindung macht und bis zu diesem Zeitpunkt keine andere Firma erfolgreich war⁶.

In den obengenannten Fixkostenmodellen ergeben sich folgende Resultate:

- Ist die Anzahl der Marktteilnehmer eine exogene Größe, nehmen die $F + E$ Ausgaben der Firma_{*i*} ab, wenn sich die Zahl der Konkurrenten erhöht.

Dies ist auch intuitiv einsichtig, wenn man sich vor Augen hält, daß mit der Zunahme der Konkurrenz der erwartete Ertrag sinkt, die Kosten hingegen konstant bleiben.

- Wird die Anzahl der Firmen endogen bestimmt, d.h. es treten so lange Firmen in den Markt, bis die erwarteten Nettoprofiten gleich Null sind, forschen, verglichen mit einer effizienten und sozial optimalen Lösung, zu viele Firmen mit zu geringen Mitteln⁷.

Hätte eine Planungsbehörde ein bestimmtes Budget für $F + E$, könnte sie die Anzahl der parallelen Forschungsprojekte bestimmen, die die erwartete Entwicklungszeit minimiert. Wären mit der „Produktionsfunktion“ $h(x)$ durchwegs zunehmende Skalenerträge verbunden, sollte bei einer effizienten Lösung das ganze Budget für ein Projekt verwendet werden, bei abnehmenden Skalenerträgen viele kleine Forschungsvorhaben und bei zunächst zunehmenden und dann abnehmenden Erträgen die Projektgröße der Bedingung $h(x_i)/x_i = h'(x_i)$ genügen, d.h. der durchschnittliche Anteil einer Geldeinheit für Forschung und Entwicklung an der Hazardrate maximal sein. Kennt eine Planungsbehörde den mit einer Innovation verbundenen gesellschaftlichen Nutzen (W^s), für jeden Zeitpunkt die Wahrscheinlichkeit für eine Innovation, wenn n Firmen mit gleich hohem Forschungsaufwand voneinander unabhängige Forschungsprojekte durchführen $P(t, x_i, n)$, ermittelt sie die optimale Anzahl der Forschungseinheiten nach folgendem Kalkül:

$$(11) \quad \max_n \rightarrow \int_0^{\infty} W^s e^{-rt} P(t, x_i, n) - nx_i.$$

Sie maximiert die Differenz aus dem Gegenwartswert des gesellschaftlichen Nutzens, der für jeden Zeitpunkt mit der Wahrscheinlichkeit seines Eintre-

⁶ $\prod_{i=1}^{n-1} [1 - P(t, x_i)] P'(t, x_n)$ gibt die Wahrscheinlichkeit für eine Innovation durch

Firma n im Zeitpunkt t bei gleichzeitiger Erfolglosigkeit der $n - 1$ Konkurrenten bis zu diesem Zeitpunkt an. Für $x_1 = x_2 = \dots = x_n$ und unter Berücksichtigung von (8) ergibt sich mit Diskontierung und Summierung über alle Zeitpunkte der Gewichtungsfaktor in (10).

⁷ In diesem Fall werden die Höhe der $F + E$ Ausgaben und die Firmenanzahl n simultan durch (10) und die Nullprofitbedingung $\int_0^{\infty} e^{-rt - nh(x)t} \pi h(x) dt = x$ bestimmt.

tens gewichtet wird, und den gesellschaftlichen Kosten über die Wahl von n , wobei die Forschungsausgaben pro Firma durch die obengenannten Effizienzüberlegungen determiniert sind.

Die Anzahl der Forschungseinheiten bei einer Marktlösung resultiert aus folgenden Überlegungen und Annahmen: Der Unternehmer erhält im Falle einer Innovation, durch Patente abgesichert, einen Wert, der dem gesellschaftlichen Nutzen einer Innovation in etwa entspricht.

Nimmt die einzelne Firma den erwarteten Zeitpunkt der Innovation als gegeben an, sieht also von ihrer Einflußmöglichkeit durch die Wahl der $F + E$ Ausgaben auf diesen Zeitpunkt ab, liegt die gewinnmaximale Ausgabenhöhe für $F + E$ im effizienten Bereich. Berücksichtigt jede Firma ihren Einfluß auf den Innovationszeitpunkt – es würde sich dann um einen Oligopolmarkt mit relativ wenigen Marktteilnehmern handeln –, investiert jede Firma etwas weniger als die effiziente Menge⁸.

Forschungsausgaben pro Firma und Firmenanzahl im Markt müssen simultan bestimmt werden, da die Ausgabenhöhe, die die Differenz aus erwartetem Ertrag und sicheren Kosten maximiert, über den Erwartungswert des Ertrages auch von der Zahl der Projekte abhängt. Zudem werden so lange Firmen in den Markt treten, bis keine Profite mehr möglich sind.

Die Firmenanzahl ist gegenüber einer Planlösung deshalb zu hoch, weil jede risikoneutrale Firma den Markt betritt, bis die erwarteten Gewinne gleich sind den Kosten. Wenn x_i bezüglich der „Technologieproduktionsfunktion“ effizient ist, wäre bei symmetrischen Nashgleichgewichten für jede Firma die Wahrscheinlichkeit im Falle einer geglückten Innovation die Firma zu sein, die erfolgreich war und damit das Patent erhält, gleich $1/n$; die durchschnittlichen Erträge sind also bei freiem Marktzutritt gleich den durchschnittlichen Kosten, und mögliche soziale Nettowohlfahrtsgewinne lösen sich auf.

Modifiziert werden die obigen Ergebnisse, wenn man nicht mehr von vertraglich fixierten Pauschalausgaben für $F + E$ im Zeitpunkt $t = 0$ ausgeht, sondern von einer „Eintrittsgebühr“ in das Patentrennen in Höhe von F und vertraglich nicht fixierte und damit wieder stochastische Ausgaben pro Zeiteinheit $x(t)$ annimmt⁹. Bis zum Innovationszeitpunkt t der Konkurrenz oder der Firma selbst stellt die Firma, bewertet zum Zeitpunkt $t = 0$,

$$(12) \quad F + \int_0^t e^{-\tau t} x ds$$

⁸ Vgl. Dasgupta / Stiglitz (1980).

⁹ Vgl. Lee / Wilde (1980).

an F + E Mitteln der Forschungsabteilung zur Verfügung. Auch in diesem Modell gibt

$$(13) \quad P(t_i \leq t) = 1 - e^{-h(x_i)t},$$

die Wahrscheinlichkeit für einen Innovationserfolg der Firma_i bis zum Zeitpunkt t an, und damit ist $h(x)$ die bedingte Wahrscheinlichkeit für eine Firma, in einem Zeitintervall $[t, t + dt]$ erfolgreich mit der Innovation auf dem Markt zu sein, wenn ihr dies bis zu diesem Zeitpunkt nicht gelungen ist. Nur hängt die Wahrscheinlichkeit jetzt von einem von der Firma zu wählenden Ausgabenstrom im Zeitablauf ab. Dieser wird durch folgendes Optimierungskalkül bestimmt¹⁰:

$$(14) \quad \max_x \rightarrow \int_0^{\infty} e^{-rt - nh(x)t} (\pi h(x) - x) dt - F.$$

Wiederum ein symmetrisches Nashgleichgewicht unterstellt, ergibt sich folgendes Resultat:

- Wird die Anzahl der Konkurrenzfirmen exogen bestimmt, steigt die Ausgabenrate für F + E jeder Firma pro Zeiteinheit, wenn die Konkurrenz zunimmt, vorausgesetzt die zu erwartenden Gewinne waren positiv.

Die dem Modell zugrundeliegende ökonomische Logik ergibt sich aus der Kostenstruktur: Bei einer Zunahme der Konkurrenz sinken nicht nur die erwarteten Gewinne, sondern in einem stärkeren Maße auch die erwarteten Kosten, da diese nicht mehr in $t = 0$ vertraglich fixiert werden und bei Marktzutritt von Firmen der Innovationszeitpunkt und damit das Zeitintervall, in dem Ausgaben getätigt werden, sich verkürzt. Ob die erwarteten Kosten stärker sinken, hängt aber von der Modellspezifikation ab.

3.0. Strategische und dynamische Innovationsaktivitäten im oligopolistisch strukturierten Markt

Im Zusammenhang mit Innovationsentscheidungen in oligopolistischen Märkten wird in der neueren Literatur häufig auf das Konzept der teilspielperfekten Gleichgewichte zurückgegriffen¹¹.

Interessant ist m.E. ein Modell, in dem die Beziehung zwischen F + E Aufwendungen und strategischen Kalkülen bezüglich der Marktanteile eines Duopols im Rahmen eines teilspielperfekten Gleichgewichtes diskutiert wird¹².

¹⁰ Die Kosten werden abdiskontiert und mit der Wahrscheinlichkeit gewichtet, daß die Innovation bis zum Zeitpunkt t nicht erfolgreich beendet wird, da nur dann weiterhin F + E Investitionen erfolgen.

¹¹ Vgl. *Selten* (1975), *Meyer / Güth* (1980), *Benoit* (1985).

Werden die $F + E$ Ausgaben nicht strategisch eingesetzt, erhält man ein einstufiges Nashgleichgewicht: Firma₁ geht von der Annahme einer fixen Outputmenge von Firma₂ aus und wählt *simultan* ihre gewinnoptimale Outputmenge und den für sie kostenminimierenden Einsatz an $F + E$ Ausgaben, wobei davon ausgegangen wird, daß die für die Produktion der Outputmenge von Firma₁ notwendigen Kosten bei Innovationstätigkeit sinken, wenn auch mit abnehmender Rate. Strategisch wird das Kalkül, wenn nicht mehr alle Kosten, inklusive der $F + E$ Kosten, für jedes Outputniveau minimiert werden, vielmehr Firma₁ den Zusammenhang zwischen der Senkung der eigenen Grenzkosten und der Erhöhung ihres Marktanteils und damit verbunden der Erweiterung ihrer Profitmöglichkeiten auf Kosten von Firma₂ in die Kalkulation aufnimmt und ihre Entscheidung bezüglich der Höhe der eingesetzten $F + E$ Mittel auf dieser neuen Grundlage trifft. Firma₁ könnte also durch eine Erhöhung der Forschungsausgaben über das kostenminimierende Niveau hinaus über die Mehreinnahmen, die aus der Reduzierung der Outputmenge von Firma₂ und der Ausweitung der eigenen Produktion resultieren, sich insgesamt besser stellen. Für Firma₂ gelten identische Überlegungen, insbesondere muß sie mit Gewinneinbußen rechnen, wenn sie die Strategien von Firma₁ ignoriert: sie wird sich ebenfalls strategisch verhalten (auf die keineswegs trivialen Stabilitätsbedingungen wird hier nicht weiter eingegangen).

Das aus dieser Konkurrenzsituation resultierende Verhalten kann als zweistufiges Spiel beschrieben werden: Auf der ersten Stufe legt jede Firma die Höhe ihrer $F + E$ Ausgaben fest und gibt diese bekannt, auf der zweiten Stufe werden die Produktionsniveaus determiniert. Teilspielperfektheit bedeutet in diesem Zusammenhang, daß durch die Wahl des $F + E$ Niveaus auf der ersten Stufe ein Mengengleichgewicht auf der zweiten Spielstufe induziert wird. Die optimalen Reaktionsfunktionen bezüglich der $F + E$ Ausgaben können nur ermittelt werden, wenn man ausgehend von der zweiten Spielstufe quasi „rückwärts“ rechnet.

Für jede Firma gilt im Optimum auf der zweiten Stufe (Mengengleichgewicht) die Gleichheit von Grenzeinnahmen in Abhängigkeit von ihrer eigenen produzierten Menge und der der Konkurrenz mit ihren Grenzkosten, die von der eigenen Produktionshöhe und den $F + E$ Ausgaben abhängen.

Aus diesen beiden Optimalitätsbedingungen erhält man für jede Firma die optimale Angebotsmenge in Abhängigkeit von den eigenen Forschungsausgaben und denen der Konkurrenzfirma. Die Gewinnfunktion für Firma_i unter Berücksichtigung des Mengengleichgewichtes auf der zweiten Stufe lautet

$$(15) \quad \pi_i = f_i(q_1(x_1, x_2); q_2(x_1, x_2); x_i) .$$

¹² Vgl. Brander / Spencer (1983 a), (1983 b).

Der Gewinn für Firma_i (π_i) ist eine Funktion des eigenen Outputs und des Outputs der Konkurrenz (q_1, q_2), wobei diese Mengen von der Höhe der Forschungsausgaben (x_1, x_2) abhängen, und wird zudem direkt vom F + E Niveau beeinflusst. Das Nashgleichgewicht bezüglich der F + E Ausgaben und die damit implizierten Reaktionsfunktionen erhält man, wenn für jede Firma die Gewinnfunktion bei gegebenem F + E Niveau der Konkurrenz nach den F + E Mitteln der Unternehmung differenziert wird und die üblichen Optimumsbedingungen erfüllt sind.

Obwohl „rückwärts“ von der zweiten Spielstufe ausgehend gerechnet wurde, ergibt sich das Nashgleichgewicht für F + E Ausgaben im Marktprozeß faktisch vor dem Mengengleichgewicht, dieses wird lediglich induziert.

Es ist daher plausibel, wenn für strategische F + E Gleichgewichte gilt:

- Im Vergleich zu den früheren Modellen investiert jede Firma bei strategischem Verhalten mehr in F + E.
- Bei symmetrischen Marktkonstellationen impliziert dies größeren Output, niedrigere Preise und weniger Profit in diesem Markt verglichen mit einem „normalen“ Duopolmarkt.

Die höheren Investitionen in F + E resultieren praktisch aus einem Gefangenendilemma, und die Mengen- und Preiseffekte sind unmittelbar einsichtig. Die Symmetrieannahme, also die Forderung nach in etwa identischen Firmen, wird mit Stabilitätsüberlegungen begründet, die hier aber nicht diskutiert werden.

Dynamische Aspekte einer F + E Konkurrenzsituation werden modelltheoretisch wie folgt berücksichtigt¹³.

Zwei in Konkurrenz zueinander stehende Firmen verfügen über ein technologisches Niveau, welches durch einen Suchvorgang aus einer gegebenen Wahrscheinlichkeitsverteilung gewonnen wurde. Mit jedem Suchvorgang sind fixe Kosten verbunden, der Planungshorizont ist unendlich.

Beide Firmen erhalten in der laufenden Periode, abhängig vom Technologieniveau der Konkurrenzfirma und dem eigenen Niveau Einnahmen und haben bei gegebenem Stand der Technologie zwei Wahlmöglichkeiten: Sie können in der laufenden Periode auf F + E verzichten und sich in späteren Perioden optimal verhalten oder aber auch als zweite Option in der laufenden Periode forschen und sich in späteren Perioden optimal verhalten.

Jede Firma wählt nun die Option, die den Gegenwartswert des erwarteten Gewinns, der natürlich auch von der gewählten Option der Konkurrenzfirma abhängt, maximiert.

¹³ Vgl. Lee (1984).

Charakteristische Eigenschaften dieser F + E Konkurrenz sind „reswitching“ und „convergence“. Während in der bisherigen Betrachtungsweise die Firma nach Optimierung ihres Erwartungsgewinnes die Strategie festlegt und nicht ändert, findet reswitching bezüglich der Optionswahl statt.

Sind z. B. beide Firmen technologisch kaum entwickelt, werden beide Firmen in F + E investieren; gelingt Firma₁ eine gegenüber Firma₂ relativ große Innovation, wird sie sich abwartend verhalten und erst nach einem größerem Erfolg von Firma₂ die Forschungstätigkeit wieder aufnehmen. Gleichgewichtslagen werden dann erreicht, wenn beide Firmen kein Interesse mehr an Forschung zeigen. Es sind auch andere ins Gleichgewicht führende Entwicklungspfade denkbar: Firma₁ könnte zunächst eine relativ große Veränderung ihrer Technologie erreichen und dann auf viele relativ kleine Technologieverbesserungen von Firma₂ nicht reagieren, um dann aber doch wieder die Option zu wechseln. Zudem konvergieren die technologischen Niveaus, da im Falle einer technologischen Lücke nun die Firma mit der geringeren technologischen Effizienz an F + E interessiert ist und daher die Differenz im technologischen Niveau sich verringert. Der genaue Spielverlauf wird einerseits davon beeinflusst, daß bei einem relativ hohen technologischen Niveau die Wahrscheinlichkeit für eine entsprechende Innovation bei konstanten Suchkosten abnimmt, andererseits die Auswirkung einer Technikänderung von Firma_i auf die Gewinne von Firma_j von der Höhe des gegebenen bzw. in Zukunft erwarteten technologischen Niveaus von Firma_i abhängen kann.

Spiel- und entscheidungstheoretische Aspekte des Innovationsprozesses lassen sich im Rahmen eines dynamischen Spieles verbinden¹⁴.

Zwei Firmen stehen im Wettbewerb um die mit einer Innovation verbundenen Gewinne, die durch ein Patentrecht über einen bestimmten Zeitraum abgesichert werden.

Jede Firma kann über den Einsatz von F + E Ressourcen sicher technisches Fachwissen über die neue Technologie akkumulieren. Die Innovation selbst ist aber stochastisch: Durch die F + E Investition und die damit implizierte Erweiterung des Fachwissens erhöht sich lediglich die Wahrscheinlichkeit eines Innovationserfolges bis zu einem bestimmten Zeitpunkt.

Die Wahrscheinlichkeit, bis zu einem bestimmten Zeitpunkt, in Abhängigkeit von dem bis dahin akkumulierten Wissen eine Innovation erfolgreich abgeschlossen zu haben, wird als objektive Wahrscheinlichkeitsverteilung interpretiert. Es wird angenommen, daß es sich um eine subjektive Wahrscheinlichkeitsverteilung handelt, die alle Spieler übereinstimmend beurteilen.

¹⁴ Vgl. *Reinganum* (1982).

Die strategische Variable ist für jede Firma die Rate, mit der sie mit Hilfe ihrer $F + E$ Ausgaben zusätzliches Wissen über die neue Technologie pro Zeiteinheit akkumuliert. Sie kann die Bestandsänderung des akkumulierten Wissens im Zeitablauf durch die Produktion neuer Informationen kontrollieren. Es handelt sich um eine sogenannte closed-loop Strategie, da die Strategie von der Zeit und dem bis zu diesem Zeitpunkt akkumulierten Wissen, auch der anderen Firma, abhängt.

Die Auszahlung für jede Firma wird von der Wahl der eigenen und der Strategie der Konkurrenz bestimmt. Dabei wird von folgendem Nashgleichgewichtskonzept ausgegangen: Jeder Spieler wählt unter der Annahme einer gegebenen Strategie der Konkurrenten die Strategie, die seine Auszahlung maximiert.

Die ökonomischen Implikationen der optimalen Strategie überraschen nicht: In jedem Zeitintervall müssen die erwarteten diskontierten Gewinne gleich sein den erwarteten diskontierten Kosten, die Forschung wird beschleunigt, wenn der Patentertrag steigt und verzögert, wenn der Planungshorizont zunimmt.

4.0. Monopolistische Kalküle bei der Bestimmung der Höhe der $F + E$ Ausgaben

Schon Schumpeter hat auf den positiven Zusammenhang zwischen Monopolmacht und Innovationsfreudigkeit hingewiesen und die Vermutung geäußert, daß große Firmen innovativer sind als kleine¹⁵.

Beide Aspekte sind voneinander unabhängig, da Monopolmacht nicht unbedingt Größe impliziert und eine große Firma nicht Monopolist sein muß.

So wird eine große Firma bei einer großen Forschungsgruppe zunächst einmal zunehmende Skalenerträge haben, andererseits kann eine zu starke Bürokratie die Leistung der Forschungsgruppe mindern.

Die monopolistische Firma kann über den Vergleich ihres gerade gegebenen Extraprofits mit dem durch eine Innovation möglichen zu dem Resultat kommen, daß sich die Innovation weniger lohnt im Vergleich zu einer Marktposition als Newcomer; andererseits kann gerade die Absicherung des laufenden Extraprofits ihre Forschungsbereitschaft erhöhen.

4.1. Die Bedeutung der Nachfrage im Innovationskalkül des Monopols

Ausgangspunkt ist eine Überlegung von *Arrow*, auf die sich später viele Autoren bezogen haben¹⁶.

¹⁵ Vgl. *Schumpeter* (1980).

Arrow untersucht den Anreiz für eine Monopolfirma, im Vergleich zu einem Wettbewerbsmarkt Ressourcen für $F + E$ aufzuwenden unter der Annahme, daß die Monopolfirma selbst die Erfindung macht, während im Wettbewerbsmarkt ein Erfinder die Innovation gegen eine Gebühr der Firma überläßt.

Die Argumentation wird mit Hilfe einer Graphik verdeutlicht.

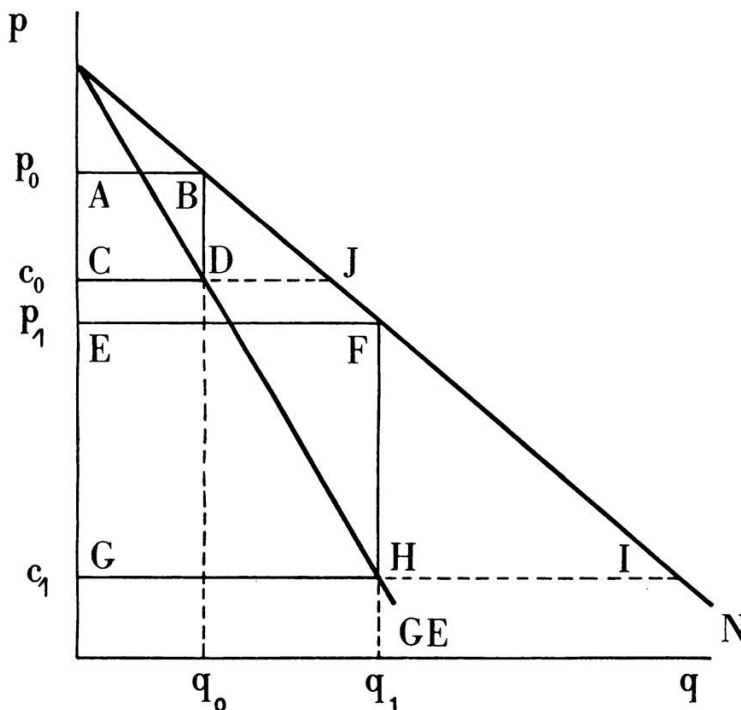


Schaubild 3

In einem Markt sei die Nachfragefunktion N gegeben und damit bei gegebenen Kosten vor einer Innovation von C_0 die gewinnmaximale Angebotsmenge (Grenzerlös = Grenzkosten) einer in diesem Markt agierenden Monopolfirma: Es wird die Menge q_0 zum Stückpreis p_0 angeboten. Der Extraprofit wird durch die Fläche $ABCD$ repräsentiert.

Senkt eine Monopolfirma durch Einsatz von $F + E$ Ressourcen die Stückkosten von c_0 auf c_1 , bietet sie die Menge q_1 zum Stückpreis p_1 an, erhält als Extraprofit $EFGH$ und damit durch eine $F + E$ Aktivität netto $EFGH - ABCD$.

¹⁶ Vgl. Arrow (1962).

Eine Innovation sei drastisch, wenn der neue Marktpreis unter den alten Grenzkosten liegt $p_1 < c_0$ und nicht drastisch, wenn $p_1 > c_0$ gilt.

Ist in dem Markt vollkommene Konkurrenz gegeben, könnte eine Forschungsfirma ihre durch ein Patent abgesicherte Innovation den Firmen gegen eine Gebühr überlassen. Es ist leicht zu sehen, daß die Forschungsfirma ihre Einnahmen bei einer Gebühr (g) pro verkaufter Einheit in Höhe von $g = p_1 - c_1$ maximiert. Bei einer nichtdrastischen Innovation könnte sie maximal $c_0 - c_1$ als Gebühr erheben.

Da die Forschungsfirma mit Nettoprofiten in Höhe von $EFGH$ rechnen kann und im Gegensatz zum Monopol nicht alte Extraprofiten ($ABCD$) verliert, ist ihr Anreiz für Innovationen größer. Überläßt die Forschungsfirma ihre Innovation Marktteilnehmern eines Oligopolmarktes, wird sie ihre Gebühr so setzen, daß sie die Differenz aus möglichen neuen und den gerade bestehenden Extraprofiten erhält. Die Gebühreneinnahmen wären am geringsten, wenn das Forschungsergebnis an eine Monopolfirma verkauft werden würde. Bei einer sozialen Planung der Innovation wäre der Forschungsanreiz am größten, da in die sozialen Gewinne noch zusätzlich Konsumentenrente eingeht. Der soziale Nutzen könnte dann in der Graphik durch die Fläche $CJIG$ repräsentiert werden.

Intuitiv einsichtig ist auch der Einfluß der Marktgröße auf den Anreiz für eine Innovation. Bei einem größeren Markt liegt die Nachfragekurve weiter rechts, der Innovator erhält mehr Gebühren, da die Technik für die Produktion von mehr Einheiten verwendet wird, ohne daß sich die Entwicklungskosten geändert haben.

Die Vorstellung, die Nachfrage und die damit verbundenen Profitmöglichkeiten seien die Hauptdeterminanten des technischen Wandels, ist durchaus kontrovers diskutiert worden. Gestützt wurde sie durch empirische Untersuchungen über die Anmeldung von Patenten¹⁷.

Dagegen wurde argumentiert, der technische Wandel müsse als zeitabhängige Variable des existierenden technologischen Niveaus begriffen werden: Wird eine Technologie einmal eingeführt, können zunächst relativ weitreichende Verbesserungen erfunden werden, im Zeitablauf nehmen die potentiellen Möglichkeiten einer Verbesserung der Technologie ab¹⁸.

Eine Synthese beider Standpunkte bietet folgender Ansatz¹⁹.

Die Nachfrage, mit der ein Monopol für seine Produkte rechnen kann, wird in zwei Komponenten zerlegt: in einen exogenen Teil, der durch Bevölkerungswachstum und Präferenzen bestimmt wird und in eine durch tech-

¹⁷ Vgl. Schmookler (1966).

¹⁸ Vgl. Kuznets (1930), Salter (1960).

¹⁹ Vgl. Gort / Wall (1986).

nologische Verbesserungen induzierten, wobei die technologische Verbesserung sich auf die Qualität des Produktes oder auf das Produktionsverfahren selbst bezieht und mit zunehmendem Reifungsgrad der Technologie kostspieliger wird.

Bei der Suche nach einem gewinnmaximalen Zeitpfad für $F + E$ Investitionen geht die Firma von einer teils exogen gegebenen zeitlichen Struktur der Nachfrage aus, berücksichtigt aber in ihrem Kalkül den Einfluß von Preissenkung und Qualitätsverbesserung auf die Nachfrage in einem bestimmten Zeitpunkt. Da Lag-Strukturen in der Optimal-Control-Theory zu Komplikationen führen, wird von einer sofortigen Nachfrageänderung bei technologischen Verbesserungen ausgegangen.

Die Nachfrage Q ist somit eine Funktion der exogenen Nachfrage D , dem Produktqualitätsindex T und der Höhe des Produktpreises.

$$(16) \quad Q = Q(D, T, P).$$

Die exogene Nachfrage wächst im Zeitablauf (S -förmig), bis eine Sättigungsgrenze erreicht wird. $T(t)$ wird im Modell bestimmt und $P = P(t)$ wird vom Monopolisten so gewählt, daß in jedem Zeitpunkt die Grenzkosten gleich sind den Grenzeinnahmen, gegeben das Technologieniveau und die exogene Nachfrage in diesem Zeitpunkt; der Preis ist also eine implizite Funktion der Technologie und der Nachfrage und muß daher in der Profitfunktion (ohne Innovationskosten)

$$(17) \quad \pi(t) = \pi[T(t), D(t)]$$

nicht explizit berücksichtigt werden.

Nach Gort und Wall ist die Änderung des Technologieniveaus im Zeitablauf $(dT/dt) = \dot{T}$ eine Funktion der $F + E$ Ausgaben pro Zeiteinheit $(x(t))$ und eine Funktion des gerade existierenden Technologieniveaus, sie spezifizieren diese Beziehung wie folgt:

$$(18) \quad \dot{T}(t) = h[T(t)] \cdot g[x(t)].$$

Das Optimierungsproblem ist damit für die Monopolfirma bekannt:

$$(19) \quad \max_x \int_0^{\infty} e^{-rt} \left[\pi(T, D) - x \right] dt$$

u. d. B. $\dot{T} = h(T)g(x); \quad T(0) = T_0 > 0, x \geq 0$

mit T als Zustands- und $x(t)$ als Kontrollvariable. Mögliche Verläufe des gewinnoptimalen Zeitpfades der $F + E$ Ausgaben $(\dot{x})$ werden unter folgenden Annahmen diskutiert:

- $\pi_T > 0$, $\pi_{TT} < 0$; die Gewinne steigen, wenn der Technologieindex steigt, jedoch mit abnehmender Rate.
- $\pi_{TD} > 0$; die marginalen Gewinne aus einer technologischen Verbesserung steigen, wenn sich die exogene Nachfrage erhöht.
- $g' > 0$, $g'' < 0$; die Zeitänderungsrate des Technologieniveaus nimmt (mit abnehmender Rate) zu, wenn die F + E Ausgaben erhöht werden.
- $h' < 0$; wenn der Reifungsgrad der Technologie zunimmt, sinkt $\dot{T}$ für gegebene F + E Ausgaben.

Explizit kann der optimale Zeitpfad für die F + E Ausgaben nur angegeben werden, wenn die Gewinn- und Technologiefunktion (T) bekannt sind. Unter relativ plausiblen Annahmen wird im Modell gezeigt, daß die Ausgaben für F + E im Zeitablauf zunächst ansteigen, ein Maximum erreichen und dann abnehmen. Die ökonomische Logik des Optimal Control Problems kann man sich folgendermaßen verdeutlichen: Würde das Monopol in einem bestimmten Zeitintervall $[t + dt]$ einen Betrag am Kapitalmarkt anlegen, kennt es genau den Wert der Ausgaben in diesem Zeitintervall und bei gegebenem Zinssatz die Einnahmen, die in der Folgezeit pro Zeitintervall generiert werden, daher auch den Gegenwartswert der zukünftigen Einnahmen in t . Wird im gleichen Zeitintervall anstatt auf dem Kapitalmarkt in F + E investiert, sind dies Kosten, denen zusätzliche Gewinne in Abhängigkeit vom Technologie- und Nachfrageniveau gegenüberstehen.

Zu Beginn der Planungsperiode, wenn die Technologie noch nicht stark entwickelt ist, sind zwar die Kosten für die Produktion einer zusätzlichen Einheit „Technologieniveau“ relativ gering, aber auch der Zuwachs an Profit hält sich in Grenzen, wenn die exogene Nachfrage ($D(t)$), deren S-förmigen Verlauf in der Planungsperiode der Unternehmer kennt, noch relativ gering ist. Es wäre also möglich, daß zu Beginn des Planungsintervalls der Ertrag aus einer F + E Investition geringer wäre als der Alternativvertrag aus einer Kapitalmarktanlage; gilt dies für den ganzen Planungszeitraum, wäre eine F + E Investition unrentabel. Da aber von einer steigenden exogenen Nachfrage ausgegangen wurde, kann zu einem späteren Zeitpunkt (wegen $\pi_{TD} > 0$) eine Investition in F + E durchaus attraktiv sein. Ist zunächst der Kapitalmarkt lukrativer, dann aber die F + E Investition, wird $\dot{x}(t) > 0$ gelten, die F + E Investitionen steigen an. Hat die exogene Nachfrage ihre Obergrenze erreicht und ist trotzdem noch ein höherer Ertrag als auf dem Kapitalmarkt zu erwirtschaften, sind die F + E Ausgaben zwar noch positiv, nehmen im Zeitablauf aber ab ($x(t) > 0$, $\dot{x}(t) < 0$), da keine zukünftigen Gewinnsteigerungen mehr antizipiert werden können.

Wenn der positive Effekt der exogenen Nachfragesteigerung auf die Gewinne ($\pi_{TD} > 0$) im Zeitablauf die Probleme einer lukrativen Technikverbesserung bei immer höherem Technologieniveau ($\pi_{TT} < 0$) überwiegt,

steigen die Grenzprofite und damit die Attraktivität einer $F + E$ Innovation gegenüber alternativen Anlagen.

Ist die Entwicklungsrate der Nachfrage relativ groß gegenüber der Abnahme der technischen Möglichkeiten und ist keine Marktsättigung zu erwarten, könnten die Ausgaben für Forschungsmittel im Zeitablauf auch permanent ansteigen. Gibt es in verschiedenen Zeitpunkten mehrere größere Innovationen, nehmen die technologischen Möglichkeiten also nicht monoton ab, wären auch mehrere Maxima für die Forschungsausgaben denkbar.

4.2. Strategische Kalkulationen einer von potentieller Konkurrenz bedrohten monopolistischen Firma im Zusammenhang mit $F + E$

Die Analyse von Arrow blendet die Konkurrenz bezüglich der $F + E$ Aktivitäten aus, die Macht eines Monopols hängt aber wesentlich von der Fähigkeit seiner Konkurrenten ab, Substitute auf den Markt zu bringen: Konkurrenz durch $F + E$ Aktivitäten ist damit für das Monopol ein wichtiges Mittel, seine Monopolsituation zu erhalten und für die Konkurrenten die Möglichkeit, ins Geschäft zu kommen.

Im folgenden Abschnitt wird der Innovationsanreiz eines Monopols im Rahmen einer Konkurrenzsituation bei Sicherheit und Unsicherheit und unter Berücksichtigung der Existenz von Lizenzmärkten diskutiert.

Die Bedingungen für einen hinreichenden Anreiz einer Innovation durch die Monopolfirma, bevor mögliche Konkurrenten auf dem Markt sind, lassen sich wie folgt entwickeln²⁰.

Die Monopolfirma produziert Gut₁, kennt für dieses Gut die Nachfragefunktion und muß mit der Patentierung von Gut₂, einem Substitut für Gut₁, durch eine neu in den Markt tretende Firma rechnen.

Der Innovationszeitpunkt T sei eine deterministische Funktion des Zeitpfades der Innovationsausgaben. Die Kosten der Innovation zum Zeitpunkt T sind somit bekannt und für alle potentiell innovativen Firmen und das Monopol identisch.

Die Strategievariablen sind für die Monopolfirmen und die möglichen Konkurrenten der Entwicklungszeitpunkt für Gut₂ und die Preisgestaltung für beide Güter, die gewinnmaximierend erfolgt. Es seien P_m^1 , P_m^2 , P_e^2 die Preise für Gut₁ und Gut₂, wie sie vom Monopol (m) bzw. der neu in den Markt tretenden Firma (e) gesetzt werden. Vor der Patentierung des Substitutes erhält das Monopol Profite mit der Rate π_m (P_m^1) pro Zeiteinheit; meldet das Monopol selbst das Patent für Gut₂ an, kann es mit einer Rate in

²⁰ Vgl. Gilbert / Newbery (1982).

Höhe von $\pi_m(P_m^1, P_m^2)$ rechnen. Im Falle einer Innovation durch eine neue Firma muß das Monopol mit einer Rate von $\pi_m(P_m^1, P_e^2)$ kalkulieren und die neue Firma erhält $\pi_e(P_m^1, P_e^2)$.

Der Innovationszeitpunkt bei Marktzutritt ist durch die Nullprofitbedingung

$$(20) \quad c(T) = \int_T^{\infty} \pi_e(P_m^1, P_e^2) e^{-rt} dt$$

determiniert. Bei freiem Zutritt zum Patentrennen kann eine Firma nicht mit Extraprofiten rechnen.

Bei Markteintritt im Zeitpunkt T wird durch

$$(21) \quad V_e = \int_0^T \pi_m(P_m^1) e^{-rt} dt + \int_T^{\infty} \pi_m(P_m^1, P_e^2) e^{-rt} dt$$

der für das Monopol mögliche Profit bei Markteintritt beschrieben. Die Monopolfirma könnte nun das Substitut nicht nur zu den Kosten $c(T)$ selbst entwickeln, sondern darüber hinaus die Entwicklung beschleunigen und Gut₂ im Zeitpunkt $T - \varepsilon$ mit dem Kostenaufwand $c(T) + d(\varepsilon)$ auf den Markt bringen. Das Monopol bleibt Monopol und kann bei „vorzeitigem“ Patentieren mit Profiten in Höhe von

$$(22) \quad V_v = \int_0^{T-\varepsilon} \pi_m(P_m^1) e^{-rt} dt + \int_{T-\varepsilon}^{\infty} \pi_m(P_m^1, P_m^2) e^{-rt} dt - \left[c(T) + d(\varepsilon) \right]$$

kalkulieren, wobei der Preis für Gut₁, den das Monopol bei vorzeitigem Patentanmeldung setzt, nicht identisch sein muß mit P_m^1 bei Markteintritt.

Die Mehreinnahmen aus vorzeitigem Patentieren ergeben sich für $\varepsilon \rightarrow 0$ gemäß²¹

$$(23) \quad V_v - V_e = \int_T^{\infty} \left[\pi_m(P_m^1, P_m^2) - (\pi_m(P_m^1, P_e^2) + \pi_e(P_m^1, P_e^2)) \right] e^{-rt} dt.$$

Der Monopolprofit bei vorzeitigem Patentieren ist dann größer als der Profit bei Markteintritt, wenn

$$(24) \quad \pi_m(P_m^1, P_m^2) > \pi_m(P_m^1, P_e^2) + \pi_e(P_m^1, P_e^2)$$

erfüllt ist. Dies ist dann der Fall, wenn durch den Marktzutritt die in dem Industriezweig insgesamt möglichen Profite durch den erhöhten Wettbewerb sinken.

²¹ Für $\varepsilon \rightarrow 0$ gilt $d \rightarrow 0$ und für $c(T)$ gilt (20).

Die Vermutung, an diesen Ergebnissen würde sich auch bei Unsicherheit nichts ändern, mußte modifiziert werden.

Es wurde auf die Bedeutung der Unsicherheit für obige Resultate im Zusammenhang mit relativ großen Innovationen hingewiesen und gezeigt, daß bei einem stochastischen Innovationsprogramm und bei hinreichend großem Marktanteil nach der Innovation, der Inkumbent (Inhaber der Monopoltechnologie) weniger in ein gegebenes Forschungsprojekt investieren wird als die Firmen, die versuchen, neu in den Markt zu kommen²². Es ist also wahrscheinlicher, daß eine herausfordernde Firma ein Patent anmeldet. Dies gilt selbst für Grenzbereiche, also gerade nicht mehr drastische Innovationen, wenn die laufenden Einnahmen hinreichend groß sind.

Die hinter diesem Modell stehende Logik ist auch intuitiv einsichtig: Investiert der Inkumbent etwas weniger in $F + E$, wird die Wahrscheinlichkeit, das Patent an den Herausforderer zu verlieren, etwas ansteigen, die Wahrscheinlichkeit, selbst das Patent zu bekommen, wird etwas sinken; gleichzeitig gibt er etwas weniger aus und sein laufender Einkommensstrom verlängert sich stochastisch etwas, weil ja die insgesamt aufgewendeten Forschungsmittel abgenommen haben und daher sich der erwartete Innovationszeitpunkt verschiebt.

Investiert die Konkurrenzfirma etwas weniger, erhöht sich die Wahrscheinlichkeit für sie, das Patent an den Inkumbenten zu verlieren, zugleich sinkt die Wahrscheinlichkeit, das Patent selbst anmelden zu können. Zwar gibt sie auch etwas weniger für $F + E$ aus, da sie aber keine laufenden Einnahmen in diesem Markt erhält, ist ihr marginaler Nutzen aus einer Reduzierung ihrer $F + E$ Aufwendungen geringer als der des Inkumbenten. Im Gleichgewicht wird diese Firma daher mehr investieren. Der marginale Nutzen aus einer Einheit zusätzlicher $F + E$ Ausgaben wird für den Inkumbenten um so geringer, je drastischer die Innovation wird: Übernimmt der Innovator nach der Innovation den ganzen Markt, würde das Monopol sich in seiner Position nur selbst ersetzen, während bei Unsicherheit durch geringere Ausgaben das Monopol durch die Verzögerung des Innovationszeitpunktes während dieser Zeitdifferenz noch Einnahmen erhält.

Zudem läßt sich zeigen, daß der Innovationsanreiz für ein Monopol im Falle von Unsicherheit um so größer wird, je kleiner (nichtdrastischer) der mit der Innovation verbundene Profit ist. Diese Argumentation kann man sich wie folgt verdeutlichen²³:

²² Vgl. *Reinganum* (1983).

²³ Vgl. v. *Ungern-Sternberg* (1980).

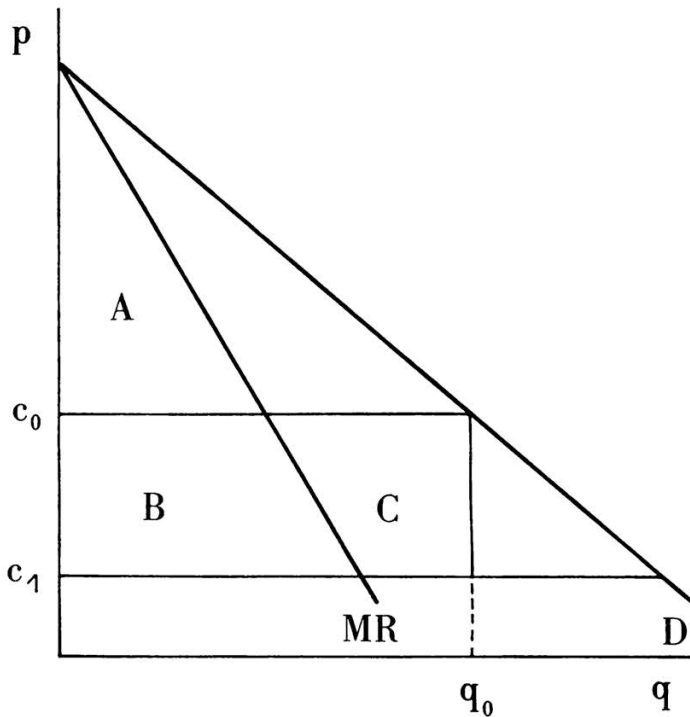


Schaubild 4

Bei gegebener linearer Nachfragekurve D und linearen konstanten Grenzkosten c_0 erhält das Monopol Profite pro Zeiteinheit, die durch die Fläche A repräsentiert werden. Entwickelt das Monopol eine Technologie, die durch konstante Grenzkosten in Höhe von c_1 gekennzeichnet ist, gibt die Fläche $A + B$ den damit verbundenen möglichen Profit pro Zeiteinheit an. Entwickelt eine neue Firma die Technologie c_1 und verhält sich nicht kooperativ, kann sie einen Preis etwas unter c_0 verlangen und die Menge q_0 absetzen. Ihr Profit pro Zeiteinheit ist dann durch $B + C$ gegeben, die Monopolprofite gehen gegen Null.

Die gewinnoptimalen $F + E$ Ausgaben des Monopols (x_m), der repräsentativen Konkurrenzfirma (x_e) und die Anzahl n der außer dem Monopol am Patentrennen teilnehmenden Firmen werden simultan durch folgende Gleichungen bestimmt²⁴:

²⁴ Dabei erhält man (25) aus dem Optimierungskalkül des Monopols

$$\max_{x_m} \int_0^{\infty} \frac{(A + B)}{r} e^{-(r + nh(x_e) + h(x_m))t} h(x_m) dt - x_m$$

$$(25) \quad \frac{(A+B) h'(x_m)}{r(h(x_m) + n h(x_e) + r)} = 1$$

$$(26) \quad \frac{(B+C) h'(x_e)}{r(h(x_m) + n h(x_e) + r)} = 1$$

$$(27) \quad \frac{(B+C) h(x_e)}{r(h(x_m) + n h(x_e) + r)} = x_e.$$

Dividiert man (26) durch (25), erhält man

$$(28) \quad h'(x_m) / h'(x_e) = (B+C) / (A+B).$$

Je größer die laufenden Profite (A) bei gegebenen (B) und (C) sind, desto größer werden die Forschungsausgaben des Monopols (x_m)²⁵.

Das Monopol hat zwei Anreize, eine Innovation durchzuführen. Einmal erhöht sich der Profit durch die Kostensenkung, und zudem sichert die monopolistische Firma ihre Stellung und die damit verbundenen Profite durch die Innovation ab.

Interessant ist m.E. auch der Hinweis auf Modifikationen des Innovationswettbewerbes, wenn die Möglichkeit einer Lizenzvergabe nach Erhalt eines Patentes eingeräumt wird²⁶. Danach wird der weniger effiziente Monopolist gegenüber den effizienteren Gegnern nicht einfach durch großen Forschungsaufwand ein frühzeitiges Patent durchsetzen, wenn die Möglichkeit exklusiver Lizenzverträge existiert, die Konkurrenten frei die Bedingungen des Lizenzvertrages aushandeln können und beide über die Auszahlungen als Monopolist oder Duopolist informiert sind. Dies kann wie folgt verdeutlicht werden:

Es seien π_m die Profite, die für eine Firma bei monopolistischer Produktion nach einem erfolgreichen Patentrennen verbleiben; bei duopolistischer Marktsituation und Symmetrie kann jede Firma mit Profiten in Höhe von π_d rechnen. Bei freiem Zutritt zum Patentrennen werden keine Extraprofite durch die Neuentwicklung gemacht, so daß die potentiell neue Firma F + E

(26) aus dem Optimierungskalkül der potentiellen Markteintreter

$$\max_{x_e} \int_0^{\infty} \frac{(B+C)}{r} e^{-(r+n h(x_e) + h(x_m)) t} h(x_e) dt - x_e$$

und (27) aus der Nullprofitbedingung. $(A+B)/r$ gibt den Gegenwartswert der Innovation im Innovationszeitpunkt an und muß daher noch diskontiert werden. Zudem wurde der Erwartungswert der Innovation $[n h(x_e) + h(x_m)]^{-1}$ von den Firmen als gegeben angesehen. Technisch bedeutet dies, daß die e -Funktion nicht nachdifferenziert werden muß.

²⁵ Vgl. Schaubild 2.

²⁶ Vgl. *Salant* (1984).

Kosten in Höhe von π_d einsetzt. Zudem wird angenommen, daß der Inkumbent nicht so effizient ist und für einen Sieg im Patentrennen Extrakosten in Höhe von d aufwenden muß.

Das Monopol wird nur dann vor der Konkurrenz ein Patent anmelden, wenn

$$(29) \quad \pi_m - (\pi_d + d) > \pi_d$$

erfüllt ist, somit der Monopolprofit gegenüber dem Duopolprofit und den Extrakosten hinreichend groß ist. Meldet eine Firma ein Patent an und bietet die Nutzung der technischen Neuerung gegen eine Gebühr in Höhe von X an, wird die Monopolfirma maximal

$$(30) \quad X = \pi_m - \pi_d$$

an Lizenzgebühr bezahlen.

Wenn sie ihr Patent nicht für X verkauft, kann sie als Monopolist produzieren, wenn sie dem Monopol dessen Patent aus der vorhergehenden Runde abkauft. Der Preis dafür wird mindestens $\pi_m - X$ betragen, und damit den Duopolprofit sicherstellen.

Um vor der neuen Firma ein Patent anmelden zu können, muß der Inkumbent mindestens $X + d$ für seine Forschung investieren. In diesem Fall erhält er $\pi_m - (X + d)$ an Profiten. Der Monopolist wird nur dann vor der Konkurrenz ein Patent anmelden, wenn

$$(31) \quad \pi_m - (X + d) \geq \pi_m - X$$

erfüllt ist. Da eine gewisse Ineffizienz des Inkumbenten angenommen wurde ($d \geq 0$), wird die monopolistische Firma im Falle der Existenz von Lizenzmärkten nie vorzeitig patentieren. Sind Monopol und neue Firma gleich effizient ($d = 0$), sind die Anreize, das Patentrennen zu gewinnen, für beide Firmen gleich groß.

Erhält die neue Firma bei den Lizenzverhandlungen um eine Geldeinheit weniger als den Betrag X , also $X - 1$, wird sie um eine Geldeinheit weniger in $F + E$ investieren und der Inkumbent muß eine Geldeinheit weniger für die vorzeitige Anmeldung eines Patentes ausgeben. Bei einer vorzeitigen Patentierung erhält der Monopolist eine Geldeinheit mehr:

$$(32) \quad \pi_m - (X - 1 + d) .$$

Damit wird aber nicht der Anreiz für eine vorzeitige Anmeldung des Patentes erhöht, da der Inkumbent auch eine Geldeinheit mehr bekommt, wenn er

das Patentrennen verliert und die Lizenz für die Produktion von der neuen Firma für $(X - 1)$ abkauft. Existieren im Zusammenhang mit den Lizenzvergaben Transaktionskosten, kann es zu einer vorzeitigen Anmeldung eines Patentes kommen, wenn die Ineffizienz des Monopolisten, gemessen in Kosten, kleiner ist als die Transaktionskosten T .

Kauft der Monopolist das Patent, bleiben ihm Profite in Höhe von $\pi_m - X - T$, verkauft er sein eigenes Patent an die neue Firma, kann er mit den gleichen Einnahmen rechnen. Die Anmeldung eines Patentes vor den Konkurrenten kostet ihn $X + d$ und nur im Falle von

$$(33) \quad \pi_m - (X + d) \geq \pi_m - X - T,$$

also wenn die Transaktionskosten größer sind als die Kosten, die durch die geringere Effizienz entstehen, wird der Monopolist vor den Konkurrenten ein Patent anmelden.

Lizenzvergabe kann ein strategisches Moment im Rahmen einer ex-ante Planung sein²⁷. Für den Monopolisten besteht die Möglichkeit durch eine Lizenzvergabe an die eventuell neu eintretende Firma, deren Anreiz eine vielleicht bessere Technologie zu entwickeln, zu mindern.

Wenn der Lizenzvertrag der neuen Firma die erwarteten Einkünfte aus weiterer Forschung sichert, wird sie auch keinen Anreiz mehr haben, ein Forschungsprojekt weiter zu verfolgen. Der Inkumbent wird das Risiko verminderter Einnahmen durch Marktzutritt gegenüber den Einnahmen aus einer Lizenzvergabe und den Kosten weiterer Forschung abwägen.

Die Auswirkungen von Imitation, Lizenzvergabemöglichkeiten und unterschiedlich striktem Patentrecht auf den F + E Wettbewerb wurden im Rahmen eines dynamischen Wartespiels analysiert²⁸.

Dabei werden die Implikationen für die F + E Investitionen untersucht, wenn die Möglichkeit für die nichtinnovative Firma existiert, durch eine Lizenzvergabe an der Innovation zu partizipieren. Ein Patentrennen wird dadurch in ein Wartespiel transformiert.

Eine Firma ist sicher an der Innovation einer anderen Firma interessiert, wenn diese ein komplementäres Gut produziert oder sie Kunde der innovativen Firma ist; kann zudem die Innovation durch Patente nicht hinreichend abgesichert werden, wird praktisch ein öffentliches Gut mit der neuen Technologie produziert und die das Innovationsrennen verlierende Firma ist daran interessiert, das Rennen so schnell wie möglich zu verlieren. Dabei wird die in der Industrie führende Firma größere Innovationen, die umfassendere Kostensenkungen implizieren, nicht durchführen, wenn das Patent-

²⁷ Vgl. Gallini (1984).

²⁸ Vgl. Katz / Shapiro (1987).

recht für die exklusive Nutzung keine hinreichende Absicherung bietet, während kleinere Innovationen unabhängig von den Imitations- und Lizenzgegebenheiten zur Ausführung kommen. Eine Senkung der Transaktionskosten für Lizenzmärkte erhöht die Profitmöglichkeit für die Lizenzgeber und bietet daher einen Anreiz für schnellere Innovationen, gleichzeitig ist es für Firmen bei Senkung der Lizenzgebühren interessanter, sich bezüglich der Innovation abwartend zu verhalten und lieber Lizenznehmer zu werden.

Abschließend sei noch auf einen Versuch hingewiesen, den Innovationswettbewerb im Zeitablauf modelltheoretisch zu erfassen²⁹. Dabei wird der „process of creative destruction“ betont, in dem eine Firma über einen bestimmten Zeitraum hinweg Monopolist bzw. Inkumbent ist, dann aber von einer neuen Firma abgelöst wird, jedoch durchaus noch die Möglichkeit hat, ihre ursprüngliche Stellung wieder einzunehmen.

Der Inkumbent ist im Markt, die herausfordernden Firmen, für die wieder Symmetrie unterstellt wird, wollen selbst in diesem Markt Monopolist werden. Da eine Sequenz von Innovationen betrachtet wird, ist durch einen einmaligen Erfolg einer herausfordernden Firma der Monopolprofit für diese Firma nicht immer sichergestellt und für den früheren Monopolisten noch nicht alles verloren. Durch einen innovativen Erfolg wird eine neue Spielstufe eröffnet, in der ein Patentrennen bis zur Erreichung einer neuen Spielstufe stattfindet. Das Spiel wird rekursiv, ausgehend von t noch verbleibenden Spielstufen entwickelt und genügt den Anforderungen eines teilspielperfekten Gleichgewichtskonzeptes. Auf jeder Spielstufe kann jede Firma die Wahrscheinlichkeit, zu einem bestimmten Zeitpunkt die Erfindung zu machen, durch die Höhe der $F + E$ Ausgaben beeinflussen. Verliert eine Firma das Patentrennen auf einer Spielstufe, hat sie die Forschungsausgaben umsonst investiert, da sie wegen der Annahme drastischer Innovationen keine Einnahmen erhält. Das nun um eine Innovation reduzierte Spiel wird aber von allen Firmen, dem neuen Inkumbenten und den neuen Konkurrenzfirmen ab dieser Stufe optimal weitergeführt.

Die eigentliche Dynamik ergibt sich in diesem Modell durch den geringeren Anreiz des Inkumbenten, in der laufenden Spielstufe in $F + E$ zu investieren wie die Herausforderer, da er in der laufenden Periode einen positiven Einkommensstrom erhält.

Da es durchaus rational für ihn ist, weniger in die Forschung zu investieren, ist es wahrscheinlich für ihn, in der nächsten Periode selbst eine herausfordernde Firma zu sein und das Spiel beginnt von neuem.

Der Wert, Inkumbent zu sein, wird größer, je weniger Spielstufen verbleiben, da zwar bei einer größeren Anzahl noch zu spielender Stufen mehr

²⁹ Vgl. *Reinganum* (1985).

Möglichkeiten für eine Niederlage des Inkumbenten existieren, aber auch die Chancen größer sind, seine Profite wiederzugewinnen.

Der Inkumbent erhöht seine Forschungsausgaben, wenn der Zinssatz oder der Wert, in der nächsten Periode Inkumbent zu sein, steigt. Seine Forschungsinvestitionen nehmen ab, wenn der laufende Profit steigt und der Wert, in der nächsten Periode Herausforderer zu sein, zunimmt.

Dieser Erklärungsansatz eines Schumpeterianischen Wettbewerbes ist m.E. der fruchtbarste, weil die Dynamik des Marktprozesses im Rahmen eines gewinnmaximierenden und spieltheoretisch konsistenten Verhaltens seitens der Firmen begründet wird.

5.0. Resümee

Im ersten Teil der Arbeit wurde die Höhe der $F + E$ Ausgaben in einem entscheidungstheoretischen Rahmen bestimmt. Die repräsentative Firma geht von einer exogen gegebenen Intensität der Konkurrenz aus und kalkuliert mit einer subjektiven Vorstellung über die Wahrscheinlichkeit einer Innovation durch die Konkurrenz bis zu einem bestimmten Zeitpunkt. Dies kann mit einer großen Anzahl konkurrierender Firmen, aber auch der Vielzahl technischer Möglichkeiten, eine Innovation zu entwickeln, begründet werden. Es wurde gezeigt, daß die Angst vor großen Verlusten, aber auch die Aussicht auf Gewinne, die Innovation beschleunigt.

Im symmetrischen Cournot-Nashgleichgewicht wird Marktunsicherheit durch technologische Unsicherheit generiert: Jede Firma kann über die Höhe der $F + E$ Ausgaben die Wahrscheinlichkeit einer erfolgreichen Innovation bis zu einem bestimmten Zeitpunkt beeinflussen. Dabei können $F + E$ Investitionen vertraglich fixiert als Pauschalausgabe zu Beginn der Planungsperiode erfolgen oder nicht vertraglich gebunden als Flußgröße pro Zeiteinheit getätigt werden. Im ersten Fall werden bei einer Verschärfung der Konkurrenzsituation die $F + E$ Ausgaben der repräsentativen Firma sinken, weil die Wahrscheinlichkeit eines Gewinnes des Patentrennens durch diese Firma und damit der erwartete Ertrag – bei weiterhin konstanten Kosten – abnimmt. Sind die Ausgaben nicht vertraglich fixiert, steigen bei einer Zunahme der Anzahl der Konkurrenzfirmen die $F + E$ Ausgaben pro Zeiteinheit, da die erwarteten Kosten wegen der stochastisch etwas kürzeren Ausgabendauer abnehmen. Nicht ganz unproblematisch dürfte in diesen Modellen die Symmetrieannahme und die Beschränkung der Betrachtungsweise auf einen Innovationszyklus sein.

Legt man dem Optimierungskalkül der Firma ein teilspielperfektes Gleichgewichtskonzept zugrunde und berücksichtigt strategisches Verhalten, ergeben sich im Vergleich zu Nashlösungen höhere gewinnoptimale $F + E$ Investitionen.

Bei nichtstochastischen Innovationen wird eine Monopolfirma „vorzeitig“ ein Patent anmelden, wenn der gerade existierende Monopolprofit die Summe der nach einer Innovation durch die Konkurrenz möglichen Duopolprofite übersteigt. Dies gilt nicht notwendigerweise, wenn die Möglichkeit einer Lizenzvergabe und stochastische Innovationen berücksichtigt werden. Für eine neue Firma im Markt ergeben sich zusätzlich zum möglichen Duopolprofit Einnahmen aus einer Lizenzvergabe, und es wurde weiterhin gezeigt, daß die Wahrscheinlichkeit einer Innovation durch den Monopolisten mit dem Umfang der Innovation abnimmt. Schließlich wurde noch auf den methodisch interessanten Versuch hingewiesen, die Schumpeterianische Vorstellung eines „process of creative destruction“ modelltheoretisch zu formulieren.

Durch die Berücksichtigung von Unsicherheit, unterschiedlichen Marktstrukturen und Verhaltensannahmen der Firmen werden neue Einsichten in die Komplexität des Innovationswettbewerbes vermittelt. Darüber hinaus wird durch die Einbeziehung des Innovationsprozesses die Evolution von Marktstrukturen verständlicher. Es wird aber auch deutlich, wie schwierig es sich gestalten dürfte, den $F + E$ Wettbewerb bei Berücksichtigung von Marktunvollkommenheiten in einem allgemeinen Gleichgewichtskonzept zu erfassen. Dies wäre insbesondere für eine wohlfahrtstheoretische Analyse wünschenswert.

Zusammenfassung

Innerhalb eines entscheidungstheoretischen Rahmens wird bei vertraglich fixierten Kosten die Innovationsgeschwindigkeit erhöht, wenn die erwarteten Einnahmen aus der Innovation steigen oder große Einbußen bei einem Verlust des Patentrennens befürchtet werden. Bei nicht vertraglich fixierten Kosten kann der Ausgabenplan revidiert werden. In einem Nashgleichgewicht nehmen die $F + E$ Ausgaben pro Firma ab, wenn die Anzahl der konkurrierenden Firmen zunimmt und die Kosten vertraglich fixiert sind; im Falle nicht vertraglich gebundener Kosten nehmen sie zu. In einem zweistufigen Duopolgleichgewichtsmodell kann gezeigt werden, daß bei strategischer Anwendung von $F + E$ der insgesamt Output und die Höhe der $F + E$ Ausgaben zunehmen. Im Falle nichtstochastischer Innovationen existiert für eine Monopolfirma ein starker Anreiz, vor möglichen Konkurrenten ein Patent anzumelden. Dies gilt nicht, wenn eine Lizenzvergabe möglich ist oder stochastische Innovationen angenommen werden. Gehen die Firmen von einer Sequenz von Innovationen aus, investiert der Inkumbent weniger in $F + E$ als die Konkurrenzfirmen.

Summary

Within the framework of decision theory, the prospect of large rewards from innovation, and the fear of heavy losses from the failure to innovate, speed up development if costs are contractual. In a Nash equilibrium in investment strategies, $R + D$ expenditure per firm declines with the number of participants in the race to innovate if costs are contractual. In a two-stage Nash duopoly model, it can be shown that strate-

gic use of R + D will increase total output and the amount of R + D expenditures. In the case of non-stochastic innovations, the monopolist has a strong incentive to patent innovations preemptively before potential entrants can do so. These results have to be modified if the licensing of patents is allowed and stochastic innovations are assumed. When firms anticipate a sequence of innovations, the industry leader invests less than a potential entrant in R + D.

Literatur

- Arrow, K. (1962), Economic Welfare and the Allocation of Ressources for Inventions, in: R. R. Nelson (Hrsg.), *The Rate and Direction of Inventive Activity*, Princeton.
- Benoit, J. (1985), Innovation and Imitation in a Duopoly. *Review of Economic Studies* 52, 98 - 106.
- Brandner, J. / Spencer, B. (1983 a), Strategic Commitment with R + D: The Symmetric Case. *Bell Journal of Economics* 14, 225 - 235.
- / — (1983 b), International R + D Rivalry and Industrial Strategy. *Review of Economic Studies* 50, 707 - 72.
- Dasgupta, P. / Stiglitz, J. (1980), Uncertainty, Industrial Structure and the Speed of R + D. *Bell Journal of Economics* 11, 1 - 28.
- Gallini, N. (1984), Deterrence by Market Sharing: A Strategic Incentive for Licensing. *American Economic Review* 74, 931 - 941.
- Gilbert, R. / Newbery, D. (1982), Preemptive Patenting and the Persistence of Monopoly. *American Economic Review* 72, 514 - 525.
- Gort, M. / Wall, R. A. (1986), The Evolution of Technologies and Investment in Innovation. *The Economic Journal* 96, 741 - 757.
- Güth, W. / Meyer, U. (1980), Innovationsentscheidungen in oligopolistischen Märkten – Eine modelltheoretische Untersuchung. *Zeitschrift für die gesamte Staatswissenschaft* 136, 113 - 115.
- Kamien, M. I. / Schwartz, N. L. (1982), *Market Structure and Innovation*. Cambridge.
- Katz, L. / Shapiro, C. (1987), R + D Rivalry with Licensing or Imitation. *American Economic Review* 77, 402 - 421.
- Kuznetz, S. (1930), *Secular Movements in Production and Prices*. New York.
- Lee, T. (1984), On the Reswitching and Convergence Properties of Research and Development Rivalry. *Management Science* 30, 186 - 197.
- Lee, T. / Wilde, L. (1980), Market Structure and Innovation: A Reformulation. *Quarterly Journal of Economics* 94, 429 - 436.
- Loury, G. C. (1979), Market Structure and Innovation. *Quarterly Journal of Economics* 93, 395 - 410.
- Reinganum, J. F. (1982), A Dynamic Game of R + D. Patent Protection and Competitive Behavior. *Econometrica* 50, 671 - 689.
- (1983), Uncertain Innovation and the Persistence of Monopoly. *American Economic Review* 73, 741 - 747.
- (1985), Innovation and Industry Evolution. *Quarterly Journal of Economics* 100, 81 - 99.

- Salant, W.* (1984), Preemptive Patenting and the Persistence of Monopoly: Comment. *American Economic Review* 74, 247 - 250.
- Salter, C. E.* (1960), Productivity and Technical Change. Cambridge.
- Schmookler, J.* (1966), Invention and Economic Growth. Cambridge, Mass.
- Schumpeter, J. A.* (1980), Kapitalismus, Sozialismus und Demokratie. 5. Aufl., Stuttgart.
- Selten, R.* (1975), Reexamination of the Perfectness Concept for Equilibrium Points in Extensive Games. *International Journal of Game Theory* 4, 25 - 55.
- Sivazlian, B. D. / Stanfel, L. E.* (1975), Analysis of Systems in Operations Research. Englewood Cliffs, N. J.
- Ungern-Sternberg v.* (1980), Current Profits and Investment Behavior. *Bell Journal of Economics* 11, 745 - 748.