



universität
wien

MASTERARBEIT

Titel der Masterarbeit

„TRANSFERPREISE FÜR IMMATERIELLE
WERTGEGENSTÄNDE“

Verfasser

Christoph Schwertl

angestrebter akademischer Grad

Master of Science (MSc)

Wien, 2013

Studienkennzahl lt. Studienblatt:
Studienrichtung lt. Studienblatt:
Betreuer / Betreuerin:

A 066 915
Masterstudium Betriebswirtschaft
Univ.-Prof. Dr. Thomas Pfeiffer

INHALTSVERZEICHNIS

TABELLENVERZEICHNIS	4
ABBILDUNGSVERZEICHNIS	4
ABKÜRZUNGSVERZEICHNIS	5
1 EINLEITUNG	6
1.1 AUSGANGSSITUATION UND PROBLEMSTELLUNG	6
1.2 ZIELSETZUNG	7
1.3 AUFBAU DER ARBEIT	7
2 TRANSFERPREIS	8
2.1 DEFINITION	8
2.2 WIRTSCHAFTLICHE ASPEKTE DES TRANSFERPREISES	8
2.3 STEUERLICHE ASPEKTE DES TRANSFERPREISES	9
3 IMMATERIELLE WERTGEGENSTÄNDE	11
3.1 DEFINITION VON IMMATERIELLEN WERTGEGENSTÄNDEN.....	11
3.2 ÖKONOMISCHE CHARAKTERISTIKA VON IMMAT. WERTGEGENSTÄNDEN	12
4 OECD GUIDELINES BASED ON ARM'S LENGTH PRINCIPLE	14
4.1 DAS ARM'S LENGTH PRINCIPLE	14
4.2 UNTERSCHIEDLICHE ARTEN VON FREMDVERGLEICHEN.....	15
4.3 TRADITIONELLE METHODEN	17
4.3.1 Die CUP-Methode	17
4.3.2 Die RP-Methode	18
4.3.3 Die CP-Methode	19
4.4 GEWINNMETHODEN	19
4.4.1 Die TNM-Methode	20
4.4.2 Die TPS-Methode.....	21
4.5 ANWENDBARKEIT DER METHODEN HINSICHTLICH IMMAT. WERTGEGENSTÄNDE	24
4.6 ANDERE BEZEICHNUNGEN IN ANDEREN LÄNDERN.....	25
4.7 ARM'S LENGTH RANGE	25
4.8 VORGEHEN ZUR TP-BESTIMMUNG	26
4.9 ABWEICHENDE AUFFASSUNGEN BEI DER TP-BESTIMMUNG	27
5 DIE DEM TP FÜR IMMAT. WERTGEGENSTÄNDE ZUGRUNDE LIEGENDE ÖKONOMIE	29

5.1	BETRACHTUNG OHNE STEUERN	30
5.2	BETRACHTUNG MIT STEUERN.....	31
5.3	REFLEXION	32
6	CUT UND CPM AUS ZENTRALER SICHT.....	33
6.1	ANALYSE DES CUT ANSATZES	33
6.1.1	Szenario 1: Ohne Division 2 entspricht First Best (FB)	34
6.1.2	Szenario 2: Mit Division 2 entspricht Second Best (SB)	35
6.1.3	Vergleich von Szenario 1 mit Szenario 2.....	36
6.1.3.1	Teilnahmebedingung greifen in beiden Szenarien.....	36
6.1.3.2	Teilnahmebedingungen greifen in beiden Szenarien nicht.....	37
6.1.3.3	Teilnahmebedingung greift in Szenario 1, aber nicht in 2	39
6.1.3.4	Teilnahmebedingung greift in Szenario 2, aber nicht in 1	40
6.1.4	Zusammenfassende Betrachtung der CUT-Analysen	40
6.2	ANALYSE DES CPM ANSATZES.....	41
6.2.1	Szenario 3: Referenzszenario	41
6.2.2	Szenario 4: CPM Szenario	43
6.2.3	Vergleich von Szenario 3 mit Szenario 4.....	45
6.2.4	Zusammenfassende Betrachtung der CPM-Analyse.....	48
6.3	REFLEXION DER GEWONNENEN ERKENNTNISSE.....	48
7	TP FÜR IMMAT. WERTGEGENSTÄNDE BEI SEQUENTIELLEM INVESTMENT.....	51
7.1	GRUNDMODELL	51
7.1.1	Operativer Profit	53
7.1.2	Transferpreis-Bestimmung	53
7.2	FIRST BEST	54
7.3	ROYALER TRANSFERPREIS BEI SEQUENTIELLEM INVESTMENT.....	54
7.3.1	Uniformer Royaler Transferpreis	54
7.3.2	Decoupled Royaler Transferpreis.....	56
7.4	VERHANDELTEN TRANSFERPREIS	58
7.4.1	Verhandelter TP kurz zusammengefasst.....	60
7.4.2	Vergleich von Verhandeltem und Royalem TP	60
7.5	ROYALER TP MIT NACHVERHANDLUNG.....	61
7.6	REFLEXION DER GEWONNENEN ERKENNTNISSE.....	63
8	DEZENTRALE PROFITMAXIMIERUNG BEI VERHANDELTEM TP ...	65
8.1	ZENTRALE UND DEZENTRALE PROFITMAXIMIERUNG.....	66
8.1.1	Zentrale Profitmaximierung - FB.....	66
8.1.2	Dezentrale Profitmaximierung mit Verhandeltem TP-SB.....	67

8.2	MANAGERBONUS ERHÖHT SICH DURCH KOOPERATION	68
8.2.1	Dezentrale Profitmaximierung mit Divisionskooperation	68
8.2.2	Probleme, welche sich durch die Divisionskooperation ergeben	69
8.3	DEZENTRALE PROFITMAXIMIERUNG DURCH DECOUPLING.....	71
8.3.1	Beschreibung der Ausgangslage	72
8.3.2	Dezentrale Profitmaximierung (SB).....	72
9	SCHLUSSKAPITEL	78
9.1	ZUSAMMENFASSUNG	78
9.2	FAZIT	79
	LITERATURVERZEICHNIS	I
	ANHANG 1	V
	ANHANG 2	VII
	ANHANG 3	IX
	ABSTRAKT	XII
	LEBENS LAUF	XIII

TABELLENVERZEICHNIS

Tabelle 1: TNM-Methode - Berechnung der Gewinnmargen	20
Tabelle 2: Aufteilung des Gewinnes nach Restgewinnmethode	23

ABBILDUNGSVERZEICHNIS

Abbildung 1: Von Boos untersuchte Modelle	29
Abbildung 2: Modell-Zeitstrahl.....	52
Abbildung 3: Investitionsniveaus bei unterschiedlichem Aufteilungsfaktor	55

ABKÜRZUNGSVERZEICHNIS

ALP	Arm's Length Principle
ALR	Arm's Length Range
CP	Cost Plus
CPM	Comparable Profit Method
CUP	Comparable Uncontrolled Price
CUT	Comparable Uncontrolled Transaction
FB	First Best
FOC	First Order Condition
MNK	Multinationaler Konzern
RP	Resale Price
SB	Second Best
TNM	Transactional Net Margin
TP	Transferpreis
TPS	Transactional Profit Split

1 EINLEITUNG

1.1 Ausgangssituation und Problemstellung

Multinationale Konzerne (MNKs) sind in der heutigen globalisierten Welt für einen großen Teil der Wirtschaftsleistung verantwortlich. Über 60% des internationalen Handels wird von MNKs durchgeführt (vgl. Markham, 2004, S. 55f).

Gleichzeitig verschieben sich die Quellen der Wertschöpfung und Wettbewerbsfähigkeit für MNKs von physischen Wertgegenständen hin zu Immateriellen Wertgegenständen (vgl. Bianchi & Labory, 2004, S. 73).

Demzufolge generieren Transferpreise für Immaterielle Wertgegenstände einen bedeutenden Teil des Profites von MNKs und wirken sich dadurch auch entscheidend auf die steuerliche Bemessungsgrundlage für die rechtlich eigenständigen Tochterunternehmen eines MNK aus.

Nationale Regierungen müssen deshalb ihre Steuersysteme entsprechend gestalten (Festsetzung des Effektivsteuersatzes), um sich beim internationalen Steuerwettbewerb mit anderen Staaten eine entsprechend vorteilhafte Position zu sichern, die dazu animiert, neue Tochterunternehmen zu gründen, bestehende zu erhalten bzw. zu erweitern und gleichzeitig maximale Steuereinnahmen für den Staat zu generieren (vgl. Ruf, 2007, S. 1ff; Ernst & Young, 2005, S. 6).

Die dadurch entstandenen Unterschiede in der Besteuerung können wiederum von MNKs benützt werden, um Ihre Gesamtsteuerlast zu senken, indem sie Gewinne¹ mittels Transferpreisen von Ländern mit einem hohen Steuersatz in Länder mit einem niedrigen Steuersatz verschieben. Besonders bieten sich dafür Immaterielle Wertgegenstände an, da diese aufgrund ihrer nicht physischen Form und schwierigen Wertbestimmung leicht zu transferieren sind (vgl. Gravelle, 2013, S. 10).

Dieses Vorgehen ist den nationalen Steuerbehörden natürlich nicht verborgen geblieben, weshalb diese die Prüfung von Transferpreisen für Immaterielle Wertgegenstände forcieren (vgl. Ernst & Young, 2009, S. 5; Mace, Faiferlick & Levy, 2007, S. 29ff).

¹ Die Begriffe Profit und Gewinn werden in dieser Arbeit als Synonym verwendet.

Berichte über die Steuersparpraktiken durch Gewinnverschiebung in Niedrigsteuerländer bekannter MNKs wie Google, Starbucks oder Amazon in der breiten Presse (vgl. Szigetvari, 2013. o. S.) rücken die Transferpreisthematik ins öffentliche Blickfeld. Diese zusätzliche Beobachtung durch die breite Bevölkerung hat insofern Früchte gezeigt, als dass Starbucks in Zukunft in Großbritannien freiwillig mehr als notwendig zahlen will (Schmid, 2012, o.S.) und gleichzeitig die OECD dazu bewogen wurde, einen neuerlichen Anlauf zu nehmen, Steuerschlupflöcher für MNKs schließen zu wollen (vgl. Szigetvari, 2013. o. S.).

Das Thema Transferpreise für Immaterielle Wertgegenstände besitzt also große aktuelle Relevanz.

1.2 Zielsetzung

Die Zielsetzung dieser Arbeit ist es, den derzeitigen Stand der Forschung hinsichtlich der Transferpreise für Immaterielle Wertgegenstände zu beleuchten. Dazu werden die entsprechenden gesetzlichen Grundlagen als Sekundärliteratur herangezogen und sämtliche Publikationen untersucht, welche sich mit ökonomischen Transferpreismodellen in Bezug auf Immaterielle Wertgegenstände beschäftigen.

1.3 Aufbau der Arbeit

Die Arbeit gliedert sich einschließlich der Einleitung in 9 Kapitel. Kapitel 2 und 3 erläutern die Theorie zu den Kernpunkten, nämlich Immat. Wertgegenstände und Transferpreise. In einem darauffolgenden Schritt werden in Kapitel 4 die gesetzlichen Rahmenbedingungen für die Transferpreisbestimmung für Immat. Wertgegenstände untersucht und an Beispielen erklärt. Aufbauend auf diese theoretischen Grundlagen werden dann anschließend in den Kapiteln 5 bis 8 vier Modelle durchleuchtet, die sich mit Transferpreisen für Immat. Wertgegenstände in unterschiedlicher Weise auseinandersetzen. Abschließend werden in Kapitel 9 die gewonnenen Erkenntnisse zusammengefasst.

2 TRANSFERPREIS

2.1 Definition

Für den Begriff Transferpreis bzw. Verrechnungspreis gibt es unterschiedliche Definitionen, wobei Verrechnungspreis und Transferpreis in dieser Arbeit als Synonyme betrachtet werden. Das Wirtschaftslexikon Gabler (vgl. Gabler Wirtschaftslexikon, 2012, o.S.) liefert die folgenden Definitionen:

- **Wirtschaftstheoretische Definition:** Preis, der nicht durch Gütertausch auf Märkten entsteht, sondern in einem Optimierungsansatz berechnet wird, auch als *Schattenpreis* bezeichnet.
- **Steuerrechtliche Definition:** Der zwischen rechtlich selbständigen, aber miteinander durch Beteiligungsbeziehungen direkt oder indirekt verbundenen Unternehmen vereinbarte Preis für Lieferungen und Leistungen jeder Art. Keine Verrechnungspreise sind folglich innerhalb einer rechtlichen Einheit, d.h. zwischen einem Stammhaus und seiner Betriebsstätte, möglich.

Das bedeutet aber nicht, dass ein und derselbe Begriff für unterschiedliche Zwecke verwendet wird, sondern dass ein Transferpreis gleichzeitig mehreren Bestimmungen gerecht werden muss.

2.2 Wirtschaftliche Aspekte des Transferpreises

Wie man aus der ersten Definition erkennen kann, bildet sich der TP (Transferpreis) nicht am freien Markt, sondern wird konzernintern festgelegt. Mit der Festsetzung des TP ist es möglich, einen konzerninternen Verkauf durchzuführen und so auch den Erfolg von Konzernteilen zu ermitteln, welche keine direkt verrechenbaren Leistungen für externe Kunden erbringen. Dadurch können diese internen Konzernteile (im weiteren Verlauf als Divisionen bezeichnet) in Konkurrenz zu externen Marktteilnehmern gesetzt werden, welche ähnliche bzw. vergleichbare Leistungen erbringen. Dabei kann es durchaus Sinn machen, den TP auf den Preis des Marktes zu setzen. Genauer dazu in Kapitel 5.

Darüber hinaus kann die Zentrale über die Festsetzung von TP Konzerneinheiten steuern (Ressourcenallokation) bzw. Anreize für ein bestimmtes Verhalten setzen.

Um diesen Sachverhalt besser veranschaulichen zu können, möchte der Autor ein kurzes Beispiel anführen. Damit dabei die Komplexität nicht unnötig erhöht wird, wird dafür ein Mat. Wertgegenstand verwendet. Der Konzern Alpha besitzt die beiden Divisionen One und Two. One produziert Zahnräder, welche von Two zur Herstellung von Zahnrادpumpen verwendet werden. Im Laufe eines Jahres liefert One Zahnräder im Wert von 100.000 € an Two. Der Einfachheit halber nehmen wir an, dass sich der Gewinn durch den TP abzüglich der anfallenden Kosten errechnet. Ob dieser TP gerechtfertigt ist oder die Zahnräder von einem Mitbewerber am freien Markt günstiger hätten bezogen werden können, soll hier nicht beurteilt werden. Dank des TP entsteht aber eine einfachere Vergleichsmöglichkeit. Gleichzeitig kann dem Geschäftsführer von One ein Anreiz zur Kostensenkung vermittelt werden, indem er einen Bonus erhält, der durch die Höhe des Gewinnes von One bestimmt wird.

2.3 Steuerliche Aspekte des Transferpreises

Zusätzlich zur internen Steuerung dient der TP auch zur Ermittlung der Steuerlast, da durch seine Hilfe der Erfolg des internen Konzernteils ermittelt werden kann.

Wenn wir das Beispiel aus dem vorherigen Abschnitt wieder aufgreifen, erkennen wir, dass der zu versteuernde Gewinn von One 100.000 € minus den dafür auftretenden Kosten ausmacht. Der Gewinn von Two hingegen verringert sich um 100.000 € gegenüber dem Gewinn, der generiert worden wäre, wenn Two keinen TP für die Zahnräder an One bezahlen hätte müssen. Wenn wir jetzt das Beispiel weiterführen und annehmen, dass One in einem Hochsteuerland und Two in einem Niedrigsteuerland ansässig ist, so erhält der Konzern Alpha durch die Festlegung der Höhe des TP für Zahnräder die Möglichkeit, seinen Konzerngewinn bestehend aus den Gewinnen von One und Two in steuerlicher Hinsicht zu optimieren, indem Alpha einen sehr niedrigen TP für die Zahnräder ansetzt und so möglichst den gesamten Gewinn im Niedrigsteuerland zu Buche schlagen lässt. Dieses Vorgehen wird aber nicht im Sinne des Geschäftsführers von One liegen (Bonus hängt vom Divisionsgewinn ab), da er ja einen möglichst hohen Bonus erhalten will.

Daraus resultiert die Erkenntnis, dass der TP mehreren Seiten gerecht werden muss/soll. Darüber hinaus wird ein niedriger TP nicht im Sinne des Hochsteuerlandes liegen, da dadurch ja seine Steuereinnahmen reduziert werden. Um das zu verhindern, gibt es für die Bildung von Transferpreisen

entsprechende gesetzliche Regelungen, welche festschreiben, wie TPs zu bilden sind. Diese Regeln werden im späteren Verlauf der Arbeit behandelt.

3 IMMATERIELLE WERTGEGENSTÄNDE

3.1 Definition von Immateriellen Wertgegenständen

Es gibt eine lange Reihe von unterschiedlichen Definitionen für den Begriff Immaterielle Wertgegenstände. Da sich die Arbeit im späteren Verlauf noch mit den OECD Guidelines für die Bestimmung von Transferpreisen befassen wird, wählt der Autor auch in diesem Abschnitt die Definition der OECD (2012, S. 7):

“..... the word “intangible” is intended to address something which is not a physical asset or a financial asset ..., and which is capable of being owned or controlled for use in commercial activities.”

Anhand dieser Definition kann man erschließen, dass es nicht sehr einfach ist, Immaterielle Wertgegenstände zu identifizieren und anschließend zu ermessen, welchen Wert bzw. welchen Vorteil sie ihren Besitzern generieren. Um diese Sache zu erleichtern, ordnet die OECD (2012, S. 9ff) Immat. Wertgegenstände in unterschiedliche Klassen ein. Darin erscheinen unter anderem:

- Patents
- Know-how and trade secrets
- Trademarks, trade names and brands
- Licences and similar limited rights in intangibles

Neben dieser Kategorisierung gibt es noch einige weitere. Die einzig für diese Arbeit relevante ist die Unterteilung in öffentliche und private Immat. Wertgegenstände (vgl. Boos, 2003, S. 18f; Lev, 2001, S. 21ff):

- Öffentliche Immat. Wertgegenstände besitzen die Charakteristika eines öffentlichen Gutes, indem sie ohne Kosten verbreitet und unbegrenzt oft eingesetzt werden können, wodurch es folglich keine Opportunitätskosten für diese Wertgegenstände geben kann. Ein Beispiel dafür wären die Rechte der Marke „McDonalds“, die weltweit bekannt ist und eine gewisse Reputation genießt. Eröffnet man ein neues Geschäft unter dieser Marke, geht ihre Reputation sofort auf das neue Geschäft über, ohne dass jene Geschäfte, die bereits unter dieser Marke betrieben werden, dadurch verlieren. Weitere Beispiele dafür wären Patente für Medikamente und Technologien.

- Private Immat. Wertgegenstände weisen genau die gegenteiligen Charakteristika wie öffentliche Immat. Wertgegenstände auf. Sie können also von der Konsumation ausgeschlossen werden. Ein Beispiel dafür wäre eine Werbekampagne für eine Hotelkette, welche aber nur auf ein spezielles Hotel zugeschnitten ist. Die anderen Hotels der Kette profitieren von dieser Kampagne deshalb nicht.

3.2 Ökonomische Charakteristika von Immat. Wertgegenständen

Wie im vorangehenden Kapitel bereits angeführt, sollte der TP hinsichtlich der Fairness gegenüber den Steuerbehörden den wirklichen Wert des transferierten Gegenstandes widerspiegeln. Das Problem dabei ist, dass der Wert von Immat. Wertgegenständen im Gegensatz zu Materiellen Wertgegenständen oft nur sehr schwer zu bestimmen ist, da es für sie keine organisierten Märkte gibt. Das liegt daran, dass Immat. Wertgegenstände besondere Charakteristika aufweisen, welche Lev (2001, S. 21ff) wie folgt beschreibt:

- Nichtrivalität: Immat. Wertgegenstände können unbegrenzt oft verwendet werden. Dadurch können keine Opportunitätskosten für Immat. Wertgegenstände entstehen.
- Höhere Skalenerträge als bei Mat. Wertgegenständen: Durch ihre Nichtrivalität können Immat. Wertgegenstände höhere Skalenerträge erreichen, da für sie keine physischen Produktionsrestriktionen gelten und aufgrund ihrer Immaterialität keine Transportkosten entstehen. Es sei denn, sie sind an einen Mat. Wertgegenstand gebunden, wie z.B. ein medizinischer Wirkstoff, der nur mittels Pillen verkauft werden kann. Dieser Zustand wird aber dadurch abgeschwächt, dass in so einem Fall die anfallenden variablen Kosten (Produktionskosten für die Pillen) im Vergleich zu den hohen Fix- bzw. Versunkenen-Kosten (Entwicklungskosten des medizinischen Wirkstoffes) vernachlässigbar sind.
- Immat. Wertgegenstände sind schwer zu managen bzw. zu steuern, da ihre Kosten und Erträge oft schwer den Produkten, Prozessen oder Aktivitäten (Activity-based costing) zuordenbar sind.
- Immat. Wertgegenstände können durch Patente nicht immer ausreichend geschützt werden. Es respektieren nämlich nicht alle

Marktteilnehmer etablierte Patente, auch wenn sie damit gegen das Gesetz verstoßen.

- Forschung & Entwicklung sowie die Etablierung erfordern hohe Investitionskosten, ohne dabei darauffolgende Erträge garantieren zu können und sind deshalb als risikoreich einzustufen (Bsp. Entwicklung eines medizinischen Wirkstoffes oder Entwicklung einer Innovation, die anschließend vom Markt nicht angenommen wird).

4 OECD GUIDELINES BASED ON ARM'S LENGTH PRINCIPLE

Nachdem die Begriffe TP und Immat. Wertgegenstand erläutert wurden, befasst sich der Autor in diesem Kapitel mit den Methoden und Regeln, welche es erlauben, den TP eines Immat. Wertgegenstandes zu bestimmen.

Regeln zur Festsetzung von TPs sind eine notwendige Komponente für jede Körperschaftssteuerregelung, da sie MNKs daran hindern können, ihre Steuerschuld zu reduzieren, indem sie ihre Profite in Niedrigsteuerländer verschieben (vgl. Brauner, 2008, S. 86).

Am häufigsten werden als Basis für solche Regelungen die OECD Guidelines on Transfer Pricing verwendet, welche erstmals 1995 publiziert wurden und seitdem fortwährend in Überarbeitung stehen. Die OECD Guidelines fungieren als gemeinsamer Rahmen für OECD Mitgliedsstaaten und bilden auch oft die Basis für Vorschriften für Entwicklungsländer, welche sich erstmals mit TPs befassen. OECD Guidelines selbst basieren auf dem Arm's length principle (ALP) und bieten unterschiedliche Vorgehensmethoden zur Festsetzung von TPs, welche im späteren Verlauf dieses Kapitels noch erörtert werden (vgl. Urquidi, 2008, S. 28f).

Zu beachten ist hierbei, dass die einzelnen Mitgliedsstaaten die OECD Guidelines nicht eins zu eins in ihre Gesetze implementieren (vgl. Markham, 2005, S. 23f). Diese Tatsache spielt in dieser Arbeit aber keine Rolle, da die einzelnen Methoden nur insoweit behandelt werden, dass der Leser deren grundlegende Funktionsweise sowie deren Anwendung in den ab Kapitel 5 beschriebenen Modellen versteht.

4.1 Das Arm's Length Principle

Das Arm's length principle (ALP) hat eine lange Geschichte und wurde das erste Mal 1935 in einem amerikanischen Gesetz formuliert (vgl. Brauner, 2008, S. 96).

Das autoritative Statement, gültig für die OECD Guidelines, findet man unter Artikel 9 der OECD Model Tax Convention und lautet wie folgt (OECDb, 2010, S. 33):

*"[Where] conditions are made or imposed between the two
[associated] enterprises in their commercial or financial*

relations which differ from those which would be made between independent enterprises, then any profits which would, but for those conditions, have accrued to one of the enterprises, but, by reason of those conditions, have not so accrued, may be included in the profits of that enterprise and taxed accordingly”.

Dieser Artikel bedeutet sinngemäß, dass eine Transaktion, welche zwischen zwei verbundenen Unternehmen stattfindet, so bewertet und damit besteuert werden muss, als würde sie zwischen nicht verbundenen Unternehmen stattfinden.

Dies lässt sich durch einen sogenannten Fremdvergleich erreichen. Dabei wird eine vergleichbare Transaktion zwischen nicht verbundenen Unternehmen als Berechnungsgrundlage genommen, um daraus einen TP zu ermitteln, welcher in steuerlicher Hinsicht für die Transaktion zwischen den beiden verbundenen Unternehmen verwendet werden kann (vgl. PWC, 2012, S. 72). Im weiteren Verlauf wird zur einfacheren Unterscheidung eine Transaktion zwischen nicht verbundenen Unternehmen als unkontrollierte Transaktion und eine Transaktion zwischen miteinander verbundenen Unternehmen als kontrollierte Transaktion bezeichnet.

Um das ALP besser anwenden zu können, wurden unterschiedliche Methoden entwickelt. Diese Methoden werden in traditionelle Methoden und in Gewinnmethoden eingeteilt, wobei dann für den jeweiligen Fall die geeignetste Methode angewendet werden soll. Dazu ist es auch möglich, mehrere Methoden gleichzeitig anzuwenden und deren Ergebnisse miteinander zu koordinieren (vgl. OECDb, 2010, S. 59ff), siehe dazu Abschnitt 4.7. Mit diesen Methoden selbst können dann noch unterschiedliche Arten von Fremdvergleichen durchgeführt werden, welche im folgenden Abschnitt beschrieben werden.

4.2 Unterschiedliche Arten von Fremdvergleichen

Zum einen kann man grundsätzlich zwischen einem inneren und einem äußeren Fremdvergleich unterscheiden.

Bei dieser Vorgehensweise kann es zur Anwendung eines inneren oder äußeren Fremdvergleichs kommen. Ein äußerer Fremdvergleich bedeutet in dieser Hinsicht, dass Transaktionen zwischen unabhängigen externen

Unternehmen betrachtet werden, ein innerer Fremdvergleich im Gegenzug betrachtet Transaktionen, welche das eigene Unternehmen mit externen Unternehmen durchgeführt hat (vgl. Verhülsdonk, 2009, S. 6f; Schreiber, 2008, S. 452).

Zum anderen kann man den Fremdvergleich noch in einen direkten und in einen indirekten Fremdvergleich unterteilen. Bei einem direkten Fremdvergleich herrscht vollständige Vergleichbarkeit (Leistung und andere wertbestimmende Faktoren wie Ressourceneinsatz und Risikoaufteilung sind bei den verglichenen Transaktionen identisch) (vgl. Verhülsdonk, 2009, S. 6f).

Dabei schreiben verschiedene Gesetzgeber unterschiedliche Qualitätsstandards vor, um unkontrollierte Transaktionen als ausreichend vergleichbar anzusehen. In Österreich z.B. müssen die folgenden 5 Vergleichsfaktoren erfüllt werden (vgl. VPR, 2010, RZ 50ff):

- Vergleichbarkeit der Produkteigenschaften
- Vergleichbarkeit der Funktion
- Vergleichbarkeit der Vertragsbedingungen
- Vergleichbarkeit der Marktgegebenheiten
- Vergleichbarkeit der Geschäftsstrategien

Auf die einzelnen Vergleichsfaktoren wird der Autor nicht näher eingehen, da diese nichts anderes als ein Leitfaden sind, um im Bedarfsfall alle Faktoren entsprechend zu bewerten und sie in die TP-Bildung einfließen zu lassen. Als solche Faktoren sind beispielsweise zu nennen: erbrachte Leistungen und übernommenes Risiko der Transaktionsteilnehmer, welche auf die Transaktion einwirken.

Diese Vergleichbarkeit ist bei einem indirekten Fremdvergleich nicht vollständig gegeben. Deshalb müssen angemessene und begründete Anpassungen durchgeführt werden, um die Vergleichbarkeit wiederherzustellen (vgl. OECDb, 2010, S. 63ff).

Aus diesem Grund ist der direkte Fremdvergleich dem indirekten Fremdvergleich vorzuziehen, weil Anpassungen immer ein Risiko mit sich bringen (vgl. PWC, 2012, S. 72f).

Die hier besprochenen Punkte gelten sowohl für die traditionellen Methoden als auch für die Gewinnmethoden und werden dort deshalb nicht mehr gesondert angeführt.

4.3 Traditionelle Methoden

4.3.1 Die CUP-Methode

Die CUP-Methode (Comparable Uncontrolled Price Method), auch als Preisvergleichsmethode bezeichnet, kommt in der Anwendung dem zuvor besprochenen ALP am nächsten und ist deshalb den anderen Methoden zum Fremdvergleich vorzuziehen, wenn alle notwendigen Voraussetzungen erfüllt werden (vgl. Verhülsdonk, 2009, S. 9). Bei der CUP-Methode werden jene Preise direkt für die kontrollierte Transaktion herangezogen, welche bei unkontrollierten Transaktionen mit vergleichbaren Produkten oder Services gehandelt werden. Je nachdem, ob eine direkte oder indirekte Vergleichbarkeit vorliegt, müssen noch entsprechende Anpassungen durchgeführt werden (vgl. OECDb, 2010, S. 63f).

Beispiel für die CUP-Methode

Zum einfacheren Verständnis wird die Grundstruktur des Beispiels aus Kapitel 2 aufgegriffen, indem der Konzern Alpha die Division One und Two vollständig kontrolliert. Division One hat ihren rechtlichen und somit steuerpflichtigen Standort in der Schweiz, während Two ebensolchen in Österreich hat. One produziert genormte Stahlzahnräder vom Typ XS und verkauft 100 Stk. um den TP an Two. Genau die gleichen Zahnräder werden von dem unabhängigen Unternehmen U1 produziert und 100 Stk. um 999 € an das ebenfalls unabhängige Unternehmen U2 verkauft. Zwischen U1 und U2 findet also eine unkontrollierte Transaktion statt und der darin erzielte Preis kann in rechtlicher Hinsicht als TP für die kontrollierte Transaktion herangezogen werden. Der zu versteuernde Gewinn von One erhöht sich deshalb um 999 € abzüglich der dafür anfallenden Kosten. Der zu versteuernde Gewinn von Two sinkt um 999 €.

Da die beiden Transaktionen vollständige Vergleichbarkeit aufweisen, sind keine Anpassungen notwendig und man spricht von einem direkten äußeren Fremdvergleich.

4.3.2 Die RP-Methode

Die RP-Methode (Resale Price Method), auch als Wiederverkaufspreismethode bezeichnet, nimmt ihren Anfangspunkt bei der konzerninternen Transaktion eines Wertgegenstandes (kontrollierte Transaktion), welcher dann später an ein nicht verbundenes Unternehmen weiterverkauft wird. Der TP nach dem ALP ergibt sich dann aus dem Preis der unkontrollierten Transaktion abzüglich der Verkaufskosten und eines angebrachten Gewinnaufschlages (Handelsspanne) (vgl. Verhülsdonk, 2009, S. 10), formell in (4.1) dargestellt.

$$TP = \text{Verkaufspreis} - (\text{Verkaufskosten} + \text{Gewinnaufschlag}) \quad (4.1)$$

Wenn keine zusätzlichen Aktivitäten durchgeführt werden und das Produkt einfach weiterverkauft wird, bildet die Gewinnspanne den fundamentalen Faktor zur Bestimmung des TP. Entsprechend dem ALP muss dieser Gewinnaufschlag durch eine vergleichbare, unkontrollierte Transaktion belegt werden können.

Komplizierter wird die RP-Methode, wenn das Produkt vor dem Weiterverkauf weiterverarbeitet bzw. verbessert wird oder es durch die Verknüpfung mit einem Immat. Wertgegenstand an Wert gewinnt. Zum Beispiel kann einem No-name-Produkt durch Hinzufügen einer Marke eine erhebliche Wertsteigerung widerfahren oder es kann für eine noch nicht bekannte Marke durch Werbung eine Werterhöhung erzielt werden. All diese Änderungen müssen wertentsprechend in die Transferpreiskalkulation eingerechnet werden (vgl. OECDb, 2010, S. 67f).

Beispiel für die RP-Methode

Für dieses Beispiel wird wieder die Grundstruktur aus dem vorherigen Abschnitt beibehalten, nur verkauft jetzt One nichtstandardisierte Spezialzahnräder an Two, welche Two dann ohne weitere Bearbeitung um 800 € pro Stk. am Markt weiterverkauft. Gleichzeitig generiert Two eine unkontrollierte Transaktion, indem sie ähnliche Spezialzahnräder um 720 € pro Stk. von U1 kauft und diese am Markt um 820 € pro Stk. wieder verkauft. Daraus errechnet sich ein Deckungsbeitrag von 100 €, der gleichzeitig die Summe aus den Verkaufskosten und dem Gewinnaufschlag darstellt. Nach der RP-Methode wird Two unterstellt, den gleichen Deckungsbeitrag auch mit den von One gelieferten Spezialzahnradern zu erzielen, wodurch auf einen TP von 700 € (800 € – 100 €) rückgerechnet werden kann.

Da die gehandelten Zahnräder sehr ähnlich sind und keine Faktoren bekannt sind, welche es notwendig erscheinen lassen bzw. es erlauben eine Anpassung durchzuführen, handelt es sich dabei somit um einen direkten inneren Fremdvergleich.

4.3.3 Die CP-Methode

Bei der CP-Methode (Cost Plus Method), auch als Kostenaufschlagsmethode bezeichnet, wird der TP für eine Transaktion zwischen verbundenen Unternehmen dadurch ermittelt, indem den Selbstkosten auf Vollkostenbasis, welche dem produzierenden Unternehmen entstehen, ein entsprechender Gewinnaufschlag hinzugerechnet wird (vgl. Verhülsdonk, 2009, S. 10; OECDb, 2010, S. 72), formell in (4.2) dargestellt:

$$TP = \text{Selbstkosten} + \text{Gewinnaufschlag} \quad (4.2)$$

Der Gewinnaufschlag wird wie bei der RP-Methode durch einen Fremdvergleich ermittelt, wobei hier auch die gleichen Regeln hinsichtlich der Wertanpassung gelten. Laut OECD eignet sich diese Methode vor allem für die Transferierung von unfertigen und halbfertigen Produkten (vgl. OECDb, 2010, S. 71f). Dem produzierenden Unternehmen wird dabei immer unterstellt, Gewinne zu erwirtschaften, ohne den Preisbildungsprozess von Angebot und Nachfrage zu berücksichtigen (vgl. Verhülsdonk, 2009, S. 11).

Beispiel für die CP-Methode

Für die CP-Methode wird das Beispiel der RP-Methode in der Hinsicht abgewandelt, dass Division Two die Spezialzahnräder nicht mehr weiterverkauft, sondern diese selbst weiterverarbeitet. Dank eines indirekten Fremdvergleichs weiß man, dass in dieser Branche bzw. für diese Produktkategorie ein Gewinnaufschlag von 10 % üblich ist. Division 1 gibt seine Selbstkosten für die Spezialzahnräder mit 650 € an. Der adäquate Gewinnaufschlag beläuft sich auf 65 € = 650 € * 0,1. Der nach dem ALP zulässige TP beträgt deshalb 715 € (650 € + 65 €).

4.4 Gewinnmethoden

Bei den unter diesem Abschnitt angeführten Methoden muss der TP nicht als konkrete Zahl, sondern kann auch als Verhältniszahl (Prozentzahl) angegeben werden.

4.4.1 Die TNM-Methode

Die TNM-Methode (Transactional Net Margin Method), auf Deutsch transaktionsbezogene Nettomargenmethode, untersucht die Nettomargen, welche bei einer kontrollierten Transaktion erzielt werden. Die Nettomarge oder auch Reingewinnspanne kann auch in Verhältnis zu einer Bezugsbasis gesetzt werden. Als Bezugsbasis können Kosten, Umsätze oder Vermögen herangezogen werden (vgl. Verhülsdonk, 2009, S. 14; OECDb, 2010, S. 77f, VPR, 2010, o.S.).

Eine Stärke der TNM-Methode im Vergleich zu den traditionellen Methoden liegt darin, dass sie durch die Verwendung der Nettomarge weniger durch transaktionelle und funktionelle Abweichungen beeinflusst wird, da funktionelle Differenzen oft auf unterschiedliche betriebliche Aufwendungen zurückzuführen sind, die der Transaktion zugeordnet werden. Andererseits gibt es allerdings auch wieder Faktoren, welche einen erheblichen Einfluss auf die Nettomarge haben, aber sich wenig auf Preis oder Bruttomarge auswirken (vgl. OECDb, 2010, S. 78f).

Beispiel für die TNM-Methode

Für die TNM-Methode nehmen wir an, dass One Autoreifenrohlinge entwickelt und diese an Two weiterverkauft (kontrollierte Transaktion). U1 produziert ebenfalls Autoreifenrohlinge und verkauft diese an U2 (unkontrollierte Transaktion). In der nachfolgenden Tabelle 1 werden alle relevanten Werte für die beiden Transaktion aus Sicht der Verkäufer dargestellt.

	One	U1
Umsatz	8.000 €	6.000 €
Herstellungskosten	3.250 €	2.500 €
Entwicklungskosten	2.500 €	2.000 €
Verwaltungskosten	1.000 €	1.000 €
Gewinn	1.250 €	500 €
Gewinnmarge	15,63%	8,33%

Quelle: Eigene Darstellung

Tabelle 1: TNM-Methode - Berechnung der Gewinnmargen

Es ist klar ersichtlich, dass die kontrollierte Transaktion zum jetzigen Zeitpunkt nicht Arm's length konform ist und ein Anpassungsbedarf besteht, weil die Gewinnmarge in Prozent gemessen am Umsatz von One fast doppelt so groß

ist wie die aus der unkontrollierten Transaktion (U1). Die Arm's length Konformität ließe sich herstellen, indem der Umsatz (Verkaufspreis) entsprechend gesenkt werden würde.

Eine völlig andere Sichtweise würde allerdings entstehen, wenn One im Gegensatz zu U1 den Reifenrohlingen zusätzlichen Wert verleiht, indem sie die Reifenrohlinge z.B. mit einer Marke versieht, die einen ausgezeichneten Ruf genießt und somit eine höhere Marge rechtfertigen kann.

Würden die Rechte für diese Marke jetzt nicht bei One, sondern bei Three (eine weitere rechtlich eigenständige Division im Konzern Alpha) liegen, so ließe sich der konforme TP dafür am einfachsten aus der Differenz der beiden Gewinnmargen errechnen. Eine diesem Vorschlag ähnliche Berechnung wird in Abschnitt 6.2 durchgeführt.

4.4.2 Die TPS-Methode

Die TPS-Methode (Transactional Profit Split Method), auf Deutsch transaktionsbezogene Gewinnaufteilungsmethode, zielt auf die Aufteilung eines durch eine Transaktion zwischen verbundenen Unternehmen generierten Gewinnes entsprechend eines Allokationsschlüssels, welcher die erbrachten Leistungen, ausgeübten Funktionen und eingegangenen Risiken der beteiligten Unternehmen widerspiegelt (vgl. Verhülsdonk, 2009, S. 14; Bakker & Obuoforibo, 2009, S. 44).

Insofern geht die TPS-Methode mit dem ALP d'accord, als dass der aufgeteilte Gewinn jenen Gewinn widerspiegelt, den nicht miteinander verbundene Unternehmen bei einer unkontrollierten Transaktion erzielt hätten. (vgl. OECDb, 2010, S. 94).

Ein Vorteil der Gewinnaufteilungsmethode ist, dass alle beteiligten Unternehmen untersucht werden. Dadurch können auch Transaktionen bewertet werden, bei denen beide Unternehmen einzigartige und wertvolle Beiträge einbringen, wie zum Beispiel Immat. Wertgegenstände (vgl. OECDb, 2010, S. 93f)

Der Vorteil, dass man ohne externe Informationen bei dieser Methode auskommt, kann im Gegenzug aber Probleme bereiten, wenn Steuerbehörden und Manager von transaktionsbeteiligten Unternehmen Umsatz und Kosten der einzelnen verbundenen Unternehmen in Erfahrung bringen wollen. Denn es ist

oft schwierig zu bestimmen, welche Kosten und Erlöse in welchem Ausmaß einer bestimmten Transaktion zugerechnet werden müssen bzw. dürfen (vgl. OECDb, 2010, S. 95).

Unterschiedliche Anwendungsformen der TPS-Methode

Zur Bestimmung des TP unter der TPS-Methode führt die OECD zwei weitere Methoden an, die Gesamtgewinnmethode und die Restgewinnmethode. Bei der Gesamtgewinnmethode wird der gesamte geschäftsfallbezogene Gewinn sofort aufgeteilt (als wären die beteiligten Unternehmen voneinander unabhängig).

Bei der Restgewinnmethode werden zuerst die Routinefunktionen der jeweiligen Division durch Anwendung der bereits behandelten Methoden abgegolten und anschließend der Restgewinn aufgeteilt. Die Aufteilung erfolgt jeweils im Verhältnis zu den eingebrachten Leistungen der Transaktionsteilnehmer (vgl. OECDb, 2010, S. 97ff).

Beispiel für die TPS-Methode

Für dieses Beispiel verwendet der Autor wieder den Konzern Alpha mit seinen beiden Division One und Two, jedoch mit der geänderten Annahme, dass jetzt beide Divisionen dieselben Produkte (Zahnräder) entwickeln, produzieren und ausschließlich an unverbundene Unternehmen verkaufen. Da beide Divisionen in der Entwicklung kooperieren und unter der gleichen Marke die Zahnräder verkaufen, folglich voneinander profitieren, wird die Gewinnaufteilungsmethode (Restwertmethode) angewandt, um den korrekten TP zu ermitteln (Tabelle 2).

Zeile		One	Two	One + Two
1	Umsatz	8.000 €	5.000 €	13.000 €
2	Herstellungskosten	3.250 €	2.000 €	5.250 €
3	Entwicklungskosten	2.500 €	1.500 €	4.000 €
4	Verwaltungskosten	1.000 €	500 €	1.500 €
5	Aufzuteilender Gewinn	1.250 €	1.000 €	2.250 €
6	Routinegewinn (10% * Herstellungskosten)	325 €	200 €	525 €
7	Aufzuteilender Restgewinn (Gewinn - Routinegewinn)			1.725 €
8	Restgewinnverteilungsschlüssel (Divisions-Entwicklungskosten / Gesamtentwicklungskosten)	62,50%	37,50%	100%
9	Restgewinn (Verteilungsschlüssel * Restgewinn)	1.078,13 €	646,88 €	1.725 €
10	Gewinn je Division nach Aufteilung	1.403,13 €	846,88 €	2.250 €

Quelle: Eigene Darstellung

Tabelle 2: Aufteilung des Gewinnes nach Restgewinnmethode

In Zeile 5 sind die Gewinne aufgeführt, welche die beiden Divisionen durch die Zahnräder erwirtschaften (Umsatz abzüglich Entwicklungs-, Verwaltungs- und Herstellungskosten). Der in Zeile 6 abgebildete Gewinn aus der Routinetätigkeit für beide Divisionen wird durch die Branchenbeobachtung (unkontrollierte Vergleichstransaktion) errechnet, sodass im Schnitt 10 % auf die Herstellkosten

aufgeschlagen werden können. Die Differenz zwischen Gesamt- und Routinegewinn ergibt den Restgewinn (Zeile 7), welcher auf Basis der Entwicklungskosten aufgeteilt wird. Zeile 8 stellt die Divisionsentwicklungskosten geteilt durch die Gesamtentwicklungskosten dar. Der dadurch entstehende Restgewinn wird in Zeile 9 angeführt. Restgewinn und Routinegewinn addiert, ergeben dann den Gesamtgewinn für die jeweilige Division. Der eigentliche TP beträgt dann 153,13 € (1000 € – 846,88 €) und überführt den Gewinn vor Aufteilung (Zeile 5) in den Gewinn nach Aufteilung (Zeile 10).

Die Aufteilung anhand der Entwicklungskosten ist Arm's length konform, da beide Divisionen die gleichen Funktionen erfüllen. Das Risiko, welches die beiden Divisionen eingehen, entspricht dem Verhältnis der Entwicklungskosten, da keine Division zusätzliche Leistungen erbringt.

4.5 Anwendbarkeit der Methoden hinsichtlich Immat. Wertgegenstände

Alle zuvor beschriebenen Methoden können zur TP Bildung für Immat. Wertgegenstände herangezogen werden. Dabei soll man immer jene Methode anwenden, welche die größtmögliche Sicherheit für die Ermittlung eines fremdvergleichskonformen TP bietet. Anders ausgedrückt heißt das, dass die Methode am zielführenden ist, bei der die wenigsten Anpassungen durchgeführt werden müssen, da Anpassungen immer ein Risiko zur Fehleinschätzung beinhalten. Bei gleicher Sicherheit sind die traditionellen Methoden den Gewinnmethoden vorzuziehen. Unter den traditionellen Methoden ist die Preisvergleichsmethode (CUP-Methode) vorrangig zu verwenden, wenn alle dafür notwendigen Gegebenheiten erfüllt sind (vgl. VPR, 2010, RZ 42ff).

Die Anwendbarkeit der traditionellen Methoden für Immat. Wertgegenstände ist allerdings äußerst begrenzt, da für sie oft passende unkontrollierte Transaktionen fehlen, die für den Vergleich herangezogen werden können. Das liegt daran, dass Immat. Wertgegenstände selten allein unkontrolliert transferiert werden. Dadurch ist nicht ersichtlich, welchen Wertbeitrag der Immat. Wertgegenstand leistet. Meist sind Immat. Wertgegenstände Teil einer ganzen Unternehmensübernahme. In den wenigen Fällen, in denen sie doch separat transferiert werden, wird normalerweise die Kompensation geheim gehalten (vgl. Brauner, 2008, S. 107; Smith & Parr, 2005, S. 169).

Die TNM-Methode umgeht das Problem, wie unter Abschnitt 4.4.1 und Abschnitt 6.2 angeführt, indem unkontrollierte Transaktionen zum Vergleich herangezogen werden, die nicht über einen Immat. Wertgegenstand verfügen und folglich der daraus resultierende Überschuss dem Immat. Wertgegenstand zugeschrieben werden kann.

Die TPS-Methode geht in ähnlicher Weise vor (siehe Abschnitt 4.4.2), da sie bei der Gesamtgewinnmethode keine, sondern nur bei der Restgewinnmethode für Immat. Wertgegenstände und Routinefunktionen unkontrollierte Vergleichstransaktionen für den Fremdvergleich benötigt.

Die beiden Gewinnmethoden brauchen also im Gegensatz zu den traditionellen Methoden auf Grund ihres Aufbaues keine Vergleichstransaktionen, die einen vergleichbaren Immat. Wertgegenstand beinhalten.

4.6 Andere Bezeichnungen in anderen Ländern

Zu bemerken ist dabei, dass verschiedene OECD Länder unterschiedliche Bezeichnungen für die beschriebenen Methoden verwenden.

Für diese Arbeit sind zwei Methoden relevant: Die amerikanische CUT (Comparable Uncontrolled Transaction) Methode, welche der OECD CUP-Methode (Comparable Uncontrolled Price) entspricht und die amerikanische CPM (Comparable Profit Method), welche der OECD TNM-Methode entspricht (vgl. PWC, 2012, S. 46).

Bei der CPM muss noch angemerkt werden, dass diese im Gegensatz zur TNM-Methode flexibler ist und verschiedene Verhältniskennzahlen erlaubt, darunter die Brutto- oder Rohgewinnspanne (vgl. Brauner, 2008, S. 131).

4.7 Arm's Length Range

Da die Bildung des TP, anders als in den vorherigen Abschnitten dargestellt, keine eindeutige Wissenschaft ist, wird es vorkommen, dass man durch die Anwendung der angemessensten Methode bzw. Methoden inklusive der notwendigen Anpassungen bei unterschiedlichen unkontrollierten Vergleichstransaktionen eine Reihe von unterschiedlichen TPs erhält (wenn kein direkter Fremdvergleich möglich ist). Wenn die notwendigen Anpassungen überhand nehmen, sollte die jeweilige Vergleichstransaktion nicht weiter berücksichtigt werden, da jede Anpassung auch ein Fehleinschätzungsrisiko mit sich bringt, welches natürlich möglichst klein gehalten werden soll. Die aus

diesem Vorgang resultierende Bandbreite an Werten wird als Arm's length range (ALR) bezeichnet. Welcher Wert innerhalb der ARL zu wählen ist, ist je nach Land unterschiedlich festgelegt. Es kann sich dabei die grundsätzliche ARL verkleinern, wenn z.B. festgeschrieben ist, dass der TP im Interquartilabstand liegen muss (vgl. OECD, 2010b, S. 2ff).

In jedem Fall sind aber für den Gesetzesprüfer entsprechende Dokumentationen und Argumentationen bereit zu halten, die die Festsetzung des TP in der entsprechenden Höhe aus gesetzlicher Sicht rechtfertigen (vgl. PWC, 2012, S. 51).

4.8 Vorgehen zur TP-Bestimmung

Prinzipiell kann zur Datenerhebung für die TP-Bestimmung jede Quelle verwendet werden. In der Praxis werden dafür verstärkt Datenbanken herangezogen. Datenbanken liefern zuverlässige Ergebnisse, wenn vergleichbare Sachverhalte untersucht werden. Falls keine vergleichbaren Transaktionen existieren (direkter Vergleich), müssen entsprechende Anpassungen durchgeführt werden (indirekter Vergleich) (vgl. Macho & Perneki, 2011, S. 294ff). Siehe dazu Abschnitt 4.2.

Welche Möglichkeiten sich daraus für Unternehmen ergeben, erkennt man, wenn man die Grundschritte eines realitätsnahen Prozesses zur TP-Ermittlung mittels der TNM-Methode betrachtet. Urquidi (2008, S. 35ff) beschreibt ihn wie folgt:

1. Im ersten Schritt wird auf elektronischem Wege, wie z.B. über die Compustat North America Datenbank von Standard & Poors, nach möglichen Vergleichsunternehmen gesucht. Anschließend wird die Suche durch weitere manuelle Recherche verfeinert, indem unpassende Unternehmen ausgeschlossen werden. Gründe für einen Ausschluss wären zum Beispiel: Das Unternehmen operiert in einer anderen geografischen Region, besitzt einzigartige Immat. Wertgegenstände (Marken, Software etc.), verfügt über ein anderes betriebliches Leistungsspektrum oder es sind zu wenige Informationen über das Unternehmen vorhanden.
2. Nach diesem Vorgehen sollte eine überschaubare Menge an Vergleichsfirmen übrig bleiben, für welche dann im zweiten Schritt die Nettomarge für die letzten Jahre berechnet wird.

3. Im dritten Schritt errechnet man dann daraus den Durchschnitt und erhält somit pro Unternehmen eine Nettogewinnmarge. Aus den Maximal- und Minimalwerten ergibt sich die ALR (falls die entsprechende Gesetzgebung keine andere Regelung vorsieht). Sie repräsentiert damit die Gewinnmarge, welche unabhängige Unternehmen für eine vergleichbare Leistung erzielen.

Der erste Schritt schafft dabei den entscheidenden Spielraum für den MNK, da hier der MNK nach eigenem Ermessen durch entsprechende Argumentation definieren kann, welche Transaktionen bzw. welche Unternehmen überhaupt für den Fremdvergleich herangezogen werden und welche Anpassungen noch notwendig sind, um eine Vergleichbarkeit herstellen zu können. Der MNK hat also einen gewissen Freiraum beim Erstellen seiner ALR. Er muss aber danach trachten, diesen im Fall der Fälle auch vor den Behörden begründen zu können (vgl. PWC, 2012, S. 74f; Brauner, 2008, S. 157).

Wie der MNK dann aus der erstellten ALR den für ihn idealen TP bestimmt, wird in Kapitel 5 und 6 behandelt.

4.9 Abweichende Auffassungen bei der TP-Bestimmung

Damit im vorangegangenen Abschnitt nicht der falsche Eindruck erweckt wird, dass die ALR von den Unternehmen mehr oder weniger frei nach eigenem Gutdünken bestimmt werden kann, sei hier gesagt, dass der TP und im Speziellen der TP für Immat. Wertgegenstände seit längerem unter besonderer Beobachtung der Behörden steht (vgl. Ernst & Young, 2003, S. 7ff; PWC, 2012, S. 17).

Um das zu untermauern, soll hier kurz der Fall von Geigy-Basle als einer von vielen erwähnt werden, bei dem die verwendeten Vergleichstransaktionen inklusive Anpassungen für die TP-Bestimmung vor Gericht entscheidend waren.

Geigy-Basle ist ein in der Schweiz ansässiges Chemieunternehmen, welches eine Division (Tochterfirma) besitzt, die den amerikanischen Markt bedient und dort auch steuerpflichtig ist. Geigy-Basle lizenziert an diese Division das Recht zwei Herbizidchemikalien zu produzieren und zu verkaufen. Im Ausgleich werden 10 % des Umsatzes an Geigy-Basle zurückgeführt. Die Steuerbehörden beklagten allerdings, dass diese 10 % an Umsatzlizenzengebühren nicht Arm's length konform waren. Die Steuerbehörden und Geigy-Basle reichten deshalb

zum Beweis ihres Standpunktes ihre Vergleichstransaktionen ein (Dokumentation, wie der ihrer Meinung nach Arm's length konformer Wert errechnet wurde). Das Gericht entschied dann, dass Geigy-Basles Vergleichstransaktionen nicht Arm's length konform waren, da in diesen Vergleichstransaktionen der Lizenznehmer die Herbizidchemikalien nur kauft und nicht selbst produziert. Dadurch muss der Lizenznehmer weniger Kapital investieren und trägt somit weniger Risiko als Geigy-Basles Division in Amerika. Allerdings stellte das Gericht auch fest, dass Geigy-Basles Herbizidchemikalien allen anderen zu jener Zeit verfügbaren Herbizidchemikalien überlegen waren. Es gewährte deshalb auf die angemessene Arm's length range einen Aufschlag, wodurch im Endeffekt wieder 10% Umsatzlizenzgebühren erreicht wurden (vgl. Brauner, 2008, S. 134ff; DeSouza, 1997. o.S.).

Weitere gerichtliche Fälle, bei denen es um den korrekten TP geht, findet man z.B. in Curd, Fickling & Minnear (2007, S. 2ff), Urquidi (2008, S. 32f); DeSouza (1997. o.S.) oder Brauner (2008, S. 136ff).

5 DIE DEM TP FÜR IMMAT. WERTGEGENSTÄNDE ZUGRUNDE LIEGENDE ÖKONOMIE

Dieser Abschnitt beschäftigt sich mit Boos (2003, S. 39ff) und den darin untersuchten Modellen, welche in Abbildung 1 dargestellt werden.

Dabei wird grundsätzlich zwischen öffentlichen und privaten Immat. Wertgegenständen (Erklärung siehe Kapitel 3) unterschieden. Des Weiteren wird differenziert, ob der Immat. Wertgegenstand mindernden Einfluss auf die Kosten hat oder die Nachfrage erhöht.

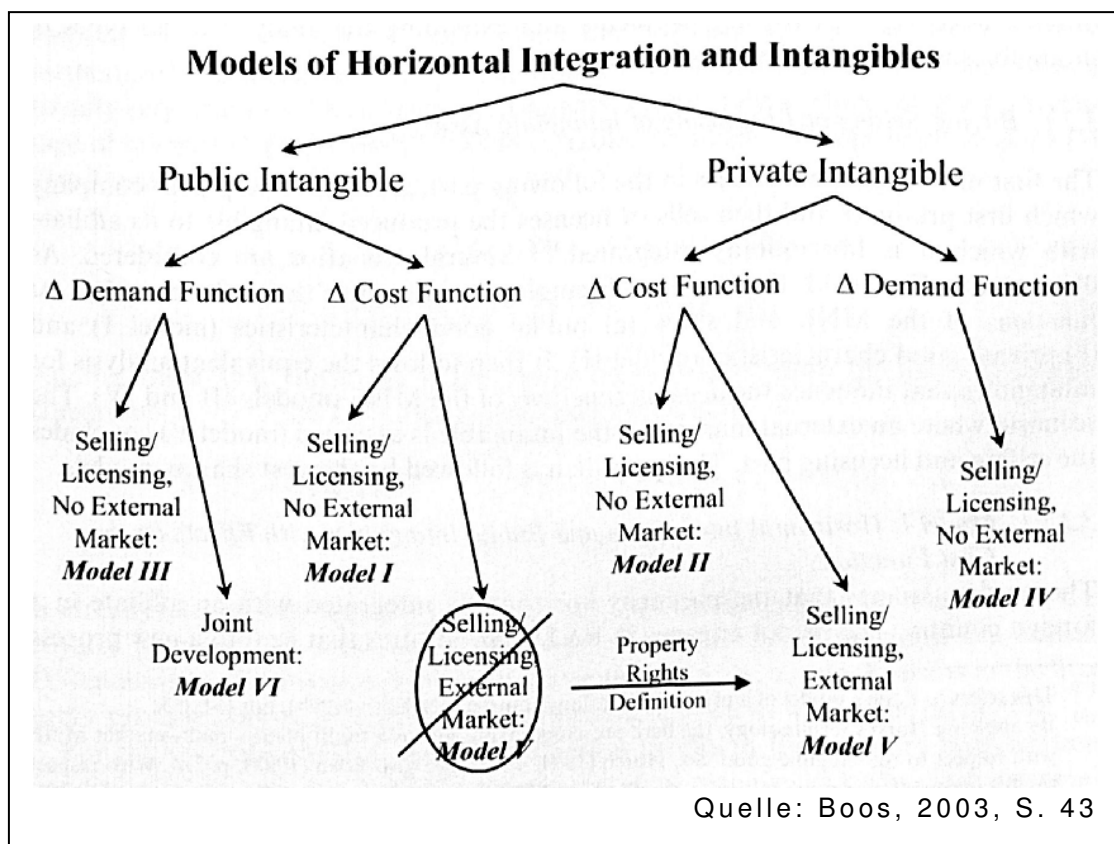


Abbildung 1: Von Boos untersuchte Modelle

Alle Modelle sind grundsätzlich so aufgebaut, dass ein MNK, der horizontal gegliedert ist, über zwei mit ihm verbundene Unternehmen, Division 1 und Division 2, verfügt. Division 1 besitzt und nutzt den Immat. Wertgegenstand. Gleichzeitig ermöglicht sie es aber, dass Division 2 ebenfalls diesen Immat. Wertgegenstand nützen kann. Diese zahlt dafür im Gegenzug einen TP an

Division 1. Die Betrachtung der Modelle erfolgt dabei immer aus Sicht der Zentrale des MNK.

Auf eine detaillierte Beschreibung der einzelnen von Boos untersuchten Modelle wird an dieser Stelle verzichtet, da deren grundsätzliche Aussagen auch ohne eine detaillierte Erklärung dargestellt werden können. Darüber hinaus werden in den nächsten Kapiteln dieser Arbeit Modelle vorgestellt, deren Aufbau und Aussagen denen von Boos (2003, S. 39ff) weitgehend entsprechen und dadurch zugleich auch die Aussagen in diesem Kapitel inhaltlich bestätigen.

5.1 Betrachtung ohne Steuern

Untersucht man diese Modelle aus Sicht der Zentrale unter dem Aspekt, den für den MNK profitmaximierenden TP zu finden, dann ergibt sich bei Außerachtlassung von Steuern, dass der durch den Immat. Wertgegenstand generierte Nutzen den durch den Immat. Wertgegenstand entstandenen Kosten entsprechen muss. Das bedeutet nichts anderes, als dass die den Immat. Wertgegenstand betreffenden Grenzkosten jenen Grenzerträgen entsprechen sollen, die durch die Nutzung desselbigen generiert werden.

Bei einem öffentlichen Immat. Wertgegenstand müssten aufgrund der Tatsache, dass durch die zusätzliche Nutzung desselben keine weiteren Kosten entstehen, alle Nutznießungen aufsummiert und den Grenzkosten des öffentlichen Immat. Wertgegenstandes gegenübergestellt werden.

Im Gegensatz dazu muss bei einem privaten Immat. Wertgegenstand nur der für die Division 2 generierte Nutzen den Grenzkosten entsprechen, da in diesem Fall genau festgestellt werden kann, welchen Vorteil Division 2 aus dem Immat. Wertgegenstand zieht. Es macht dabei keinen Unterschied, ob die Nutzung des Immat. Wertgegenstandes kostensenkend oder nachfragesteigernd wirkt.

Es ist aber sinnvoll, bei einem kostensenkenden Immat. Wertgegenstand den TP auf Basis der Kostensenkung zu berechnen. Eine Berechnung auf Basis der Nachfrage wäre nicht sinnvoll, da sich darin nicht der Einfluss des Immat. Wertgegenstandes widerspiegeln würde. Vice versa gilt dasselbe für einen nachfragesteigernden Immat. Wertgegenstand.

Im Modell V wird noch gesondert untersucht, welchen Einfluss ein wettbewerbsintensiver externer Markt auf die Bildung des TP für einen privaten Immat. Wertgegenstand hat. Das daraus resultierende Ergebnis ist auch als Hirshleifer-Regel bekannt und besagt, dass der MNK seinen Profit maximiert, indem er den TP auf Höhe des Marktpreises festsetzt, wobei dieser Preis dann auch dem Grenznutzen entsprechen muss, welcher durch den Immat. Wertgegenstand generiert wird, sowie jenen Grenzkosten, welche für Division 1 entstehen. Falls dieser Preis über dem optimalen TP (bei rein interner Betrachtung) liegt, muss dieser auf den Marktpreis gesenkt werden, um den Profit zu maximieren.

5.2 Betrachtung mit Steuern

Inkludiert man Steuern auf den Gewinn in die vorherigen Betrachtungen, so erhält man dieselben Ergebnisse wie vorher, jedoch unter rechnerischer Berücksichtigung der Steuern.

Der durch den Immat. Wertgegenstand generierte Nutzen nach Steuern muss also den durch den Immat. Wertgegenstand entstandenen Kosten nach Steuern entsprechen. Das bedeutete, dass die den Immat. Wertgegenstand betreffenden Grenzkosten nach Steuern jenen Grenzerträgen nach Steuern entsprechen sollen, die durch die Nutzung des Immat. Wertgegenstandes generiert werden.

Nimmt man dabei an, dass beide Divisionen unterschiedliche Steuersätze auf den Divisionsgewinn zahlen müssen, da sie in verschiedenen Ländern angesiedelt sind, so kann der Gesamtprofit des MNK optimiert werden bzw. die Gesamtsteuerlast gemindert werden, indem der TP zur Verschiebung der Gewinne verwendet wird. Wenn der Steuersatz auf den Gewinn von Division 1 höher ist als auf den Gewinn von Division 2, so wird der TP so hoch wie möglich (obere Grenze der ALR, siehe Abschnitt 4.7) angesetzt werden. Ist der Steuersatz auf den Gewinn von Division 1 niedriger als von Division 2, so wird der TP so niedrig wie möglich (untere Grenze der ALR) angesetzt werden. Voraussetzung dafür ist, dass es zu keiner Doppelbesteuerung kommt.

Eine anschaulichere Betrachtung dieses Sachverhaltes unter Berücksichtigung unterschiedlicher Rahmenbedingungen erfolgt im anschließenden Kapitel 6 bzw. in Kapitel 8.

5.3 Reflexion

In dieser vom Autor behandelten Publikation von Boos (2003) wird aufgezeigt, dass der TP für Immat. Wertgegenstände aus zentraler Sicht unter Berücksichtigung der speziellen Eigenschaften von öffentlichen und privaten Immat. Wertgegenständen zur Optimierung des Gesamtprofites genauso gewählt werden muss, als ob es sich um Materielle Wertgegenstände handeln würde. Steuern haben darauf grundsätzlich keinen Einfluss. Allerdings können Steuerdifferenzen durch den TP ausgenützt werden, um den Gesamtprofit des MNK zu heben bzw. die Gesamtsteuerlast zu senken.

Darüber hinaus hat das Modell V, welches den TP bei Vorhandensein eines externen Marktes für Immat. Wertgegenstände behandelt, so gut wie keine Relevanz, da solche Märkte nicht vorhanden sind (vgl. Gu & Lev, 2001, S. 2).

Eine Betrachtung der Modelle bei dezentraler Profitmaximierung wurde von Boos nicht durchgeführt. Dies wird in den Kapiteln 7 und 8 in dieser Arbeit nachgeholt.

6 CUT UND CPM AUS ZENTRALER SICHT

In diesem Kapitel werden die im Journalartikel "U.S. Income Tax Transfer Pricing Rules for Intangibles as Approximations of Arm's Length Pricing" von Halperin und Srinidhi (1996, S. 61 – 80) vorgestellten Modelle zur Transferpreisermittlung mittels CUT und CPM (OECD CUP- und TNM-Methode) untersucht. Das Hauptaugenmerk liegt dabei darauf, inwieweit es multinationalen Konzernen möglich ist, die Transferpreisermittlung zu beeinflussen.

Dabei wird angenommen, dass alle Entscheidungen von der Zentrale gefällt werden. Dadurch kann ausgeschlossen werden, dass es zu wirtschaftlichen Koordinationskonflikten zwischen einzelnen Divisionen innerhalb des MNK kommen kann.

6.1 Analyse des CUT Ansatzes

Die amerikanische CUT (Comparable Uncontrolled Transactions) Methode entspricht grundsätzlich der OECD CUP (Comparable Uncontrolled Price) Methode. Der TP wird deshalb in diesem Modell, wie in Abschnitt 4.3.1 beschrieben, von einer vergleichbaren unkontrollierten Transaktion übernommen.

In diesem Modell wird angenommen, dass ein MNK (Zentrale) eine Tochterfirma (Division 1) im Land U (Ursprungsland des Immat. Wertgegenstandes) besitzt, welche somit dem Steuersatz t_U ($T_U = 1 - t_U$) unterliegt und eine Tochterfirma (Division 2) im Land F, die dort dem Steuersatz t_F ($T_F = 1 - t_F$) ausgesetzt ist.

Division 1 stellt Division 2 einen Immat. Wertgegenstand zur Verfügung, was Division 2 in die Lage bringt, die inverse Nachfragefunktion $(K * B - q_M)$ für das von ihr produzierte Produkt M mit der Stückzahl q_M um den Faktor K weiter nach rechts parallel zu verschieben. Es wird angenommen, dass der Immat. Wertgegenstand die Charakteristika eines öffentlichen Gutes aufweist. Es fallen deshalb keine Transaktions- und Opportunitätskosten an. Ein Beispiel für solch einen Immat. Wertgegenstand wäre das Recht, eine bereits etablierte und bekannte Marke verwenden zu dürfen. Das bedeutet, dass für Division 2 durch die Verwendung eines aus dem Land U stammenden Immat. Wertgegenstandes Vorteile entstehen, wenn $K > 1$ ist, diese allerdings nicht im Land U versteuert.

Da das gegen die gesetzlichen Vorschriften verstößt, bezahlt Division 2 den Betrag r pro verkauftem q_M als TP für die Nutzung des Immat. Wertgegenstandes. Der Wert von r wird dabei aus einer unkontrollierten vergleichbaren Transaktion übernommen, was dem Arm's length Prinzip entspricht.

Diese unkontrollierte Vergleichstransaktion entsteht in diesem Modell durch die Lizenzierung des gleichen Immat. Wertgegenstandes um den Betrag r pro q_I an ein nicht mit dem MNK verbundenes Unternehmen I, welches im Land G operiert, das einen Steuersatz t_G ($T_G = 1 - t_G$) aufweist. Der Zukauf des Immat. Wertgegenstandes ermöglicht I die inverse Nachfragefunktion ($K * B - D * q_I$) für sein Produkt q_I um den Faktor K weiter nach rechts parallel zu verschieben.

Der Wert D , wie in (6.1) dargestellt, zeigt das Verhältnis der Marktgrößen zwischen dem Markt, auf dem Division 2 operiert und jenem Markt, auf dem der unverbundene Lizenznehmer I agiert.

$$D = \frac{\text{Nachfrage in F bei gegebenem Preis (p)}}{\text{Nachfrage in G bei gegebenem Preis (p)}} \quad (6.1)$$

Teilnahmebedingung

Es zeigt sich aus Sicht des Lizenznehmers I, dass die durch den Immat. Wertgegenstand verursachten Kosten durch den Mehrwert, den er generiert, ausgeglichen werden. Was nichts anders heißt, als dass der Profit von I mit dem Immat. Wertgegenstand größer/gleich dem Profit ohne Immat. Wertgegenstand sein muss (eine ausführliche Herleitung bieten Halperin & Srinidhi (1996, S. 75)). Daraus ergibt sich folgende Teilnahmebedingung $r \leq (K - 1) * B$, die erfüllt werden muss, damit die Verwendung des Immat. Wertgegenstandes für I als sinnvoll zu erachten ist.

Sowohl für Division 2 als auch für das nichtverbundene Unternehmen I (Lizenznehmer) werden die variablen Kosten pro q mit v und die Fixkosten mit f bezeichnet. Die Produktionskosten betragen bei beiden Unternehmen q^2 .

6.1.1 Szenario 1: Ohne Division 2 entspricht First Best (FB)

Im ersten zur Untersuchung stehenden Szenario ist die Division 2 des MNK nicht existent. Die Division 1 des MNK maximiert ihren Profit (6.2) durch die

Optimierung von r , wobei in diesem Szenario r mit r_0 bezeichnet wird, da dieser Wert den fairen Wert hinsichtlich der Transferpreisbestimmungen darstellt. Der Profit nach Steuern ergibt sich dabei nur aus den versteuerten Lizeinahmen. Der Lizenznehmer I wiederum berechnet seinen Profit (6.3) durch q_I , multipliziert mit seiner inversen Nachfragefunktion abzüglich der Produktionskosten q_I^2 , der Fixkosten f , der variablen Kosten v pro q_I und der Lizenzierungskosten r_0 pro q_I .

$$\Pi = \max_{q_I, r_0} \{ T_u * q_I * r_0 \} \quad (6.2)$$

Unter folgenden Bedingungen:

$$q_I = \max_{q_I} \{ T_G[q_I(K * B - D * q_I) - q_I^2 - v * q_I - q_I * r_0 - f] \} \quad (6.3)$$

$$r_0 \leq (K - 1) * B \quad (6.4)$$

Im ersten Schritt maximiert der Lizenznehmer seinen Profit, indem er das optimale q_I (6.3) bestimmt: durch Bilden der FOC hinsichtlich q_I und Umformung des Ergebnisses auf q_I . Das so errechnete optimale q_I setzt man in (6.2) ein, bildet die FOC nach r_0 und formt sie anschließend auf r_0 um. Das ergibt folgendes r_0 (6.5) bei Nichtgreifen der Teilnahmebedingung (6.4).

$$r_0 = (K * B - v) * 0,5 \quad (6.5)$$

Falls (6.4) greift, hat r_0 den Wert aus (6.6),

$$r = (K - 1) * B \quad (6.6)$$

6.1.2 Szenario 2: Mit Division 2 entspricht Second Best (SB)

Dieses Szenario nimmt das ursprüngliche Modell wieder auf bzw. erweitert Szenario 1 durch Hinzufügen von Division 2 des MNK.

Die Profitfunktion des MNK (6.7) beinhaltet deshalb im ersten Teil die Einnahmen von Division 1 aus der Lizenzierung an Division 2 und an das nichtverbundene Unternehmen I (Lizenznehmer), versteuert mit dem Steuersatz t_U ($T_U = 1 - t_U$). Der zweite Teil ergibt sich aus der Profitfunktion

von Division 2. Diese besteht aus der inversen Nachfragefunktion, multipliziert mit q_M abzüglich der Produktionskosten q_M^2 , der Fixkosten f , der variablen Kosten v pro q_M und der Lizenzierungskosten r pro q_M .

Im Gegensatz dazu hat sich die Profitfunktion des Lizenznehmers I gegenüber Szenario 1 nicht verändert.

$$\Pi = \max_{q_M, q_I, r_0} \{ T_U * r(q_M + q_I) * T_F[q_M(K * B - q_M) - q_M^2 - q_M(v + r) - f] \} \quad (6.7)$$

Unter den Bedingungen:

$$q_I = \max_{q_I} \{ T_G[q_I(K * B - D * q_I) - q_I^2 - v * q_I - q_I * r - f] \} \quad (6.8)$$

$$r \leq (K - 1) * B \quad (6.9)$$

Nachdem der Lizenznehmer seinen Profit durch Bestimmung des optimalen q_I berechnet hat (6.8), wird (6.7) partiell nach q_M abgeleitet und schließlich die so bestimmten q_I und q_M in (6.7) erneut eingesetzt, die FOC nach r gebildet und nach r umgeformt. Das ergibt folgendes r (6.10) bei Nichtgreifen von (6.9):

$$r = (K * B - v) \left[\frac{2 * T_U * T_F - (L + 1) T_F (T_F - T_U)}{4 * T_U * T_F - (L + 1) (T_F - T_U)^2} \right] \quad (6.10)$$

Falls (6.9) greift, hat r den Wert aus (6.11):

$$r = (K - 1) * B \quad (6.11)$$

6.1.3 Vergleich von Szenario 1 mit Szenario 2

6.1.3.1 Teilnahmebedingung greifen in beiden Szenarien

Analysiert man nun den Preis für den Immat. Wertgegenstand r_0 (6.6) bzw. r (6.11) aus den beiden Szenarien, so ist klar ersichtlich, dass bei Greifen der Teilnahmebedingungen (6.4) und (6.9) r_0 gleich r ist (TP in FB und SB gleich). Der optimale TP für den MNK ist also gleich hoch, wenn zusätzlich zur

Lizenzierung an den Lizenznehmer I auch Division 2 den Immat. Wertgegenstand als Lizenznehmer verwendet.

6.1.3.2 Teilnahmebedingungen greifen in beiden Szenarien nicht

Bei diesem Zustand ändert der MNK den Wert von r_0 aus (6.5) auf den Wert für r aus (6.10). Dabei müssen die unterschiedlichen Relationen der Steuersätze zueinander in Betracht gezogen werden.

Fall 1: $T_U = T_F$

Wenn die Steuersätze auf den Profit in beiden Ländern gleich hoch sind ($T_U = T_F$), dann ist r aus (6.10) gleich r_0 aus (6.5). Der optimale Preis r zur Profitmaximierung in SB entspricht also dem FB-Preis.

Fall 2: $T_U < T_F$

In diesem Fall ist der Steuersatz im Land U höher als im Land von Division 2 ($t_U > t_F$). Daraus resultiert, dass bei der Berechnung von r in (6.10) der Term in den eckigen Klammern kleiner als 0,5 ist, was bedeutet, dass r_0 größer als r ist. Der MNK senkt nun den Preis für den Immat. Wertgegenstand, wodurch sich die Einnahmen aus der Lizenzierung an das unverbundene Unternehmen I verringern. Im Gegenzug kann jedoch die Gesamtsteuerbelastung des MNK durch Verringerung des TP r gesenkt werden, da so weniger des von Division 2 erzielten Gewinns (niedrigerer TP) in U versteuert wird, aber dafür mehr in F.

Die durch die Preissenkung entstandenen Verluste auf Grund der geringeren Einnahmen aus der Lizenzierung an I in SB gegenüber FB unter Berücksichtigung der eingerechneten Mengensteigerung durch die Senkung der Grenzkosten, werden durch die Steuerersparnis mehr als ausgeglichen, welche durch die Gewinnverschiebung erreicht wird.

Es ist dabei zu berücksichtigen, dass durch die Transferpreissenkung und der damit einhergehenden Steuerreduktion q_M steigt. Aufgrund der Steuersenkung fallen auch die Grenzkosten (Steuern sind ein Teil der Kosten), wodurch die optimale Stückzahl (q_M) gegenüber Szenario 1 steigt.

Inwieweit r unter r_0 angesetzt wird, hängt zum einen mit der Steuerdifferenz ($T_F - T_U$) zusammen, da mit zunehmender Differenz der Preis r sinkt (Partielle Ableitung von (6.10) nach T_F ist negativ). Mit zunehmendem T_F sinkt r , gleichzeitig erhöht sich dadurch die Steuerdifferenz. Zum anderen sinkt r mit

wachsendem D (Partielle Ableitung von (6.10) nach D ist negativ). Das kann soweit führen bis r gleich null ist. Das tritt ein, wenn (6.12) erfüllt wird. (6.12) errechnet man, indem man (6.10) gleich null setzt und auf D umformt.

$$D \geq \frac{3 * T_U - T_F}{T_F - T_U} \quad (6.12)$$

Je größer der Markt von Division 2 im Vergleich zum Markt des Lizenznehmers ist, umso mehr Profit lässt sich aus der Preisreduzierung und somit aus der Steuerverminderung schlagen.

Fall 3: $T_U > T_F$

Wenn der Steuersatz im Land U geringer als im Land F ($t_U < t_F$) $\rightarrow (T_U > T_F)$ ist, dann wird der MNK natürlich bestrebt sein, möglichst wenig Profit in F zu versteuern und den Profit stattdessen nach U zu verschieben, was sich durch einen möglichst hohen TP erreichen lässt. Es tritt also genau die Umkehrung des zuvor besprochenen Falles 2 auf. Der Wert r wird größer als der Wert r_0 angesetzt.

Dabei ist aber zum einen zu beachten, dass die Erhöhung des TP nicht jenen Profit erhöht, welcher aus dem Lizenzverkauf an das nichtverbundene Unternehmen I resultiert, da dieser durch r_0 maximiert wird.

Zum anderen kann r maximal bis zur bestehenden Erfüllung von (6.9) $r \leq (K - 1) * B$ erhöht werden, denn bis zu diesem Betrag lohnt es sich für das nichtverbundene Unternehmen I den Immat. Wertgegenstand zuzukaufen. Wenn sich der Zukauf nicht mehr lohnt, also die (6.9) nicht mehr erfüllt wird, bricht das ganze Modell zusammen.

Inwieweit r über r_0 angesetzt wird, hängt zum einen mit der Steuerdifferenz ($T_F - T_U$) zusammen, da mit zunehmender Differenz der Preis r steigt (Partielle Ableitung von (6.10) nach T_F ist positiv). Mit zunehmenden T_F steigt r , gleichzeitig erhöht sich dadurch die Steuerdifferenz. Zum anderen steigt r mit wachsendem D (Partielle Ableitung von (6.10) nach D ist positiv).

Je größer der Markt, in dem Division 2 operiert, gegenüber dem Markt des Lizenznehmers I ist, umso mehr Profit lässt sich aus der Preissteigerung und somit aus der Steuerverminderung schlagen.

6.1.3.3 Teilnahmebedingung greift in Szenario 1, aber nicht in 2

Bei diesem Zustand greift die Teilnahmebedingung in Szenario 1. Das bedeutet, dass der maximale Preis erreicht ist, den der Lizenznehmer I bereit ist für den gegebenen Immat. Wertgegenstand zu bezahlen. r kann also nicht mehr wachsen, da sonst die für das Modell notwendige unkontrollierte Vergleichstransaktion verloren geht.

Sobald aber der Immat. Wertgegenstand zusätzlich auch noch der Division 2 zur Verfügung gestellt wird, kann der Preis sinken oder gleich bleiben. Sinkt der TP, so ist für den Lizenznehmer die Teilnahmebedingung nicht mehr bindend.

Fall 1: $T_U = T_F$

Zur Erinnerung sei gesagt: Wenn in Szenario 1 (FB) die Teilnahmebedingung greift, gilt $r_0 = ((K - 1) * B)$ (6.6). Bei Nichtgreifen der Teilnahmebedingung in Szenario 2 weist r den Wert von (6.10) auf. Setzt man in (6.10) $T_U = T_F$ ein, ergibt sich $r = (K * B - v) * 0,5$ (Das entspricht r_0 bei Nichtgreifen der Teilnahmebedingung).

Da $r < r_0$ ist, wird der TP in SB gegenüber FB gesenkt und so der Profit für den MNK optimiert. Interessant ist, dass dieser Sachverhalt auch dann seine Gültigkeit behält, wenn D sehr klein ist, d.h. der Markt des Lizenznehmers sehr viel größer als der Markt von Division 2 ist. Es erscheint unlogisch, dass unter diesen Rahmenbedingungen (Markt von D viel größer als der von Division 2) der MNK den Preis für den Immat. Wertgegenstand senkt, da sich in so einem Fall mehr Profit aus der Lizenzierung an I erzielen lässt als aus der Lizenzierung an Division 2.

Aus Sicht des Autors ist dieser Sachverhalt auf eine nicht perfekte Implementierung von D zurückzuführen. D kann nämlich nur zur Senkung des TP führen (ersichtlich aus (6.10)).

Fall 2: $T_U < T_F$

Dieser Fall bringt genau dieselben Erkenntnisse wie Fall 2 in Abschnitt 6.1.3.2, allerdings bei leicht veränderter Ausgangsstellung, denn r_0 ist bei Greifen der Teilnahmebedingung größer als bei Nichtgreifen der Teilnahmebedingung. Das bedeutet, dass die Differenz zwischen r_0 und r größer ist. Die Ergebnisse gleichen aber denen in Abschnitt 6.1.3.2 Fall 2.

Fall 3: $T_U > T_F$

Dieser Fall entspricht im Grunde Fall 3 aus Abschnitt 6.1.3.2, denn r hat den gleichen Wert und weist somit grundsätzlich auch die gleichen Ergebnisse auf. Allerdings ist r kleiner als r_0 , da r_0 bereits den maximalen Preis aufweist, den der Lizenznehmer bereit ist zu zahlen. Der Wert r strebt also mit wachsender Steuerdifferenz auf den Wert r_0 zu, was gegenteilig zu Fall 3, Abschnitt 6.1.3.2 ist. Der Rest entspricht Abschnitt 6.1.3.2 Fall 3.

6.1.3.4 Teilnahmebedingung greift in Szenario 2, aber nicht in 1

Bei diesem Zustand greift die Teilnahmebedingung in Szenario 1 nicht, jedoch in Szenario 2. Das bedeutet, dass der MNK in der SB-Betrachtung den TP auf den maximal möglichen Wert erhöht, den der Lizenznehmer noch bereit ist für die Benutzung des Immat. Wertgegenstandes zu zahlen (entspricht Teilnahmebedingung).

Dass der MNK in SB den TP erhöht, wie in den Abschnitten 6.1.3.2 und 6.1.3.3 ersichtlich ist, kann für die folgenden Fälle aber nicht gelten: ($T_U = T_F$) und ($T_U < T_F$), da in diesen Fällen der TP nicht steigen darf.

Wenn bei diesem Zustand (Teilnahmebedingung greift in Szenario 2, aber nicht in Szenario 1) der Fall eintritt, dass ($T_U > T_F$) ist, trifft der in Fall 3, Abschnitt 6.1.3.3 beschriebene Extremfall (r erreicht $[(K - 1) * B]$) zu.

6.1.4 Zusammenfassende Betrachtung der CUT-Analysen

Wie die unterschiedlichen Analysen gezeigt haben, ist der MNK durch Veränderung des TP passend zu steuerlichen Rahmenbedingung in der Lage, seinen Profit zu erhöhen. Die CUT-Methode bietet dem MNK also eine Möglichkeit, unter passenden Umständen seine Gesamtsteuerbelastung durch Verschiebung der steuerpflichtigen Gewinne in das Land mit dem niedrigen Steuersatz zu senken.

Für alle Analysen gilt, dass zwischen den Ländern U und F ein Abkommen besteht, das eine Doppelbesteuerung verhindert. Wäre das nicht der Fall, hätte der MNK immer Anlass gehabt, den TP zu senken, da der TP in beiden Ländern zur Steuerbemessungsgrundlage beiträgt.

6.2 Analyse des CPM Ansatzes

Die amerikanische CPM (Comparable Profit Method) Methode entspricht der OECD TNM-Methode (Transactional Net Margin Method), siehe Abschnitt 4.4.1.

Zur Analyse des CPM Ansatzes wird das Modell, welches beim CUT Ansatz verwendet wurde, in der Hinsicht abgeändert, dass der Immat. Wertgegenstand nur mehr der Division 2 zur Verfügung steht. Gleichzeitig agiert das nichtverbundene Unternehmen, welches um Komplikationen zu vermeiden jetzt mit L (für lokales Unternehmen stehend) bezeichnet wird, auf dem gleichen Markt wie Division 2. Auf diesem Markt bilden die beiden Firmen ein Duopol, wobei Division 2 als Stackelberg-Führer agiert. Die inverse Nachfragefunktion des lokalen Unternehmens lautet deshalb $(1 * B - q_m - q_L)$. Im Vergleich dazu ermöglicht der Immat. Wertgegenstand der Division 2 folgende inverse Nachfragefunktion $(K * B - q_m - q_L)$. Die anderen Variablen sind gleich wie beim CUT Ansatz (Abschnitt 6.1).

Der Teil des Profits des MNK, welcher in U versteuert werden muss und somit den Wert des Immat. Wertgegenstandes widerspiegelt, wird in zwei Schritten ermittelt. Im ersten Schritt wird der Profit des lokalen Unternehmens L auf Division 2 umgemünzt, indem man die Kennzahl Rohertrag (Umsatz- minus Produktionskosten) zu betrieblichen Aufwendungen (variable Kosten plus Fixkosten) (6.25) des lokalen Unternehmens L mit der Kennzahl Rohertrag zu betrieblichen Aufwendungen von Division 2 gleichsetzt und daraus den Profit errechnet, welchen Division 2 ohne Immat. Wertgegenstand erzielt hätte. Anschließend (zweiter Schritt) wird die Differenz dieses errechneten Profits zum tatsächlich erzielten Profit von Division 2 im Land U und der restliche Profit im Land F versteuert.

Im weiteren Verlauf der Arbeit werden wieder zwei Szenarien behandelt, wobei beim ersten (Szenario 3: Referenzszenario) angenommen wird, dass die Steuerbehörden volle Information über den Wert des Immat. Wertgegenstandes besitzen. Im darauffolgenden Szenario (Szenario 4) besitzen die Steuerbehörden diese Information nicht mehr. Der Wert des Immat. Wertgegenstandes wird deshalb durch die CPM Methode ermittelt.

6.2.1 Szenario 3: Referenzszenario

In diesem Referenzszenario erhält man die korrekte Information über den Wert des Immat. Wertgegenstandes durch Vergleich der Profitmaximierung ohne

(erster Teil) und mit (zweiter Teil) dem Immat. Wertgegenstand. Die Differenz zeigt dann den Wert des Immat. Wertgegenstandes.

Erster Teil: Szenario 3:

Im ersten Teil dieses Szenarios wird angenommen, dass auch Division 2 über keinen Immat. Wertgegenstand verfügt. Daraus ergibt sich folgendes Optimierungsproblem:

$$\Pi_Z^{OI} = \max_{q_{L31} \quad q_{M31}} \{ T_F * \Pi_{M31}^{OI} \} \quad (6.13)$$

Unter den Bedingungen:

$$\Pi_{M31}^{OI} = [q_{M31}(1 * B - q_{M31} - q_{L31}) - q_{M31}^2 - v * q_{M31} - f] \quad (6.14)$$

$$q_{L31} = \max_{q_L} \{ [q_L(B - q_{M31} - q_L) - q_L^2 - v * q_L - f] \} \quad (6.15)$$

(6.15) versichert, dass das lokale Unternehmen seinen Profit maximiert. Division 2 inkludiert die daraus entstehende optimale Menge q_{L31} in (6.14) und kann so ihren Profit (6.13) durch Bestimmen der optimalen Menge q_{M31} maximieren.

Zweiter Teil: Szenario 3:

In diesem Teil des Szenarios besitzt Division 2 wieder den Immat. Wertgegenstand, wobei dessen Wertbeitrag, welcher sich aus der Profitdifferenz zwischen Szenario 3 -Teil 1 und 2 ergibt, vollständig im Land U $t_U(T_U = 1 - t_U)$ versteuert wird. Der restliche Wertbeitrag wird im Land F versteuert $t_F(T_F = 1 - t_F)$, wie in (6.16) ersichtlich ist.

$$\Pi_Z = \max_{q_L \quad q_M} \{ T_F * \Pi_{M31}^{OI} + T_U(\Pi_M - \Pi_{M31}^{OI}) \} \quad (6.16)$$

Unter den Bedingungen:

$$\Pi_M = [q_M(K * B - q_M - q_L) - q_M^2 - v * q_M - f] \quad (6.17)$$

$$q_L = \max_{q_L} \{ [q_L(B - q_M - q_L) - q_L^2 - v * q_L - f] \} \quad (6.18)$$

Die Bedingungen (6.17) und (6.18) funktionieren in gleicher Weise wie die Bedingungen (6.14) und (6.15) in Szenario 3 / Teil 1 und erfüllen auch denselben Zweck. Die sich daraus ergebenden maximierenden Mengen lauten:

$$q_L = \frac{B - q_M - v}{4} \quad (6.19)$$

$$q_M = \frac{2}{7} \left(\frac{B(4 * K - 1)}{4} - \frac{3 * v}{4} \right) \quad (6.20)$$

Diese q_L und q_M werden produziert, wenn die Steuerbehörden den konkreten Wert des Immat. Wertgegenstandes kennen und somit der TP den wirklichen Wert des Immat. Wertgegenstandes darstellt. Dieser wird dadurch auch korrekt versteuert.

6.2.2 Szenario 4: CPM Szenario

In diesem Szenario kennen nun die Steuerbehörden den Wert des Immat. Wertgegenstandes nicht. Zur Steuerermittlung wird deshalb mittels der CPM Methode der Profit des lokalen Unternehmens L auf den Profit von Division 2 umgelegt (siehe (6.25)), weil L nicht über den Immat. Wertgegenstand verfügt. Dies geschieht mit Hilfe der Kennzahl Rohertrag (Umsatz- minus Produktionskosten) zu betrieblichen Aufwendungen (variable Kosten plus Fixkosten). Zur Errechnung des Rohertrages werden deshalb zum Profit die zuvor schon abgezogenen variablen Kosten plus Fixkosten wieder addiert. Dieser umgelegte Profit stellt den Profit ohne Immat. Wertgegenstand (Π_{M2}) für Division 2 dar. Da angenommen wird, dass zwischen dem Land F und dem Land U ein Abkommen existiert, welches eine Doppelbesteuerung verhindert, maximiert sich der Profit des MNK durch (6.21). Dabei wird der Profit ohne Immat. Wertgegenstand (Π_{M2}) nur in F versteuert und der wirkliche Profit von Division 2 (Π_M) abzüglich Π_{M2} , im Land U versteuert, wobei $\Pi_M - \Pi_{M2}$ den Wert des Immat. Wertgegenstandes widerspiegelt.

$$\Pi_Z = \max_{q_L, q_M} \{ T_F * \Pi_{M2} + T_U(\Pi_M - \Pi_{M2}) \} \quad (6.21)$$

Unter den Bedingungen:

$$\Pi_M = [q_M(K * B - q_M - q_L) - q_M^2 - v * q_M - f] \quad (6.22)$$

$$\Pi_L = [q_L(B - q_M - q_L) - q_L^2 - v * q_L - f] \quad (6.23)$$

$$q_L = \max_{q_L} \{ \Pi_L \} \quad (6.24)$$

$$\frac{\Pi_{M2} + f + v * q_M}{f + v * q_M} = \frac{\Pi_L + f + v * q_L}{f + v * q_L} \quad (6.25)$$

Zuerst wird durch (6.24) das maximierende q_L bestimmt. Dieses q_L wird dann in die restlichen Bedingungen (6.22), (6.23) und (6.25) eingesetzt. Diese werden dann wiederum in die Profitfunktion (6.21) eingesetzt. Die Zentrale maximiert jetzt ihren Profit, indem sie das optimale Produktionsniveau von q_M bestimmt.

Daraus ergeben sich die folgenden Werte für q_L und q_M . Um den anschließenden Vergleich zu vereinfachen, berechnen wir die Mengen einmal unter reiner Fixkostenbetrachtung ($v = 0$)

$$q_L = \frac{B - q_M}{4} \quad (6.26)$$

$$q_M = \frac{-(T_F - T_U) \frac{B}{4} + T_U(B(K - \frac{1}{4}))}{-(T_F - T_U) \frac{1}{4} + T_U \frac{7}{2}} \quad (6.27)$$

und einmal unter reiner Betrachtung der variablen Kosten ($f = 0$).

$$q_L = \frac{B - q_M - v}{4} \quad (6.28)$$

$$q_M = \frac{(T_F - T_U) \frac{B - v}{2} + T_U(B(K - \frac{1}{4}) - \frac{3 * v}{4})}{(T_F - T_U) + \frac{7}{2} T_U} \quad (6.29)$$

6.2.3 Vergleich von Szenario 3 mit Szenario 4

In Szenario 4 hängen die Mengen von den Steuersätzen in den jeweiligen Ländern ab, deshalb werden die Mengen unter Annahme unterschiedlicher Steuervoraussetzungen analysiert.

Fall 1: $T_U = T_F$

Bei Betrachtung der q_M aus (6.27) bzw. (6.29) und (6.20) sollte ersichtlich sein, dass bei gleicher Steuerbelastung ($T_U = T_F$) die Referenzmengen den Mengen unter CPM entsprechen.

$$q_M (6.20) = q_M (6.29)$$

$$\frac{2}{7} \left(\frac{B(4 * K - 1)}{4} - \frac{3 * v}{4} \right) = \frac{(T_F - T_U) \frac{B - v}{2} + T_U (B \left(K - \frac{1}{4} \right) - \frac{3v}{4})}{(T_F - T_U) + \frac{7}{2} T_U}$$

Das gilt natürlich auch bei reiner Fixkostenbetrachtung, $q_M (6.20) = q_M (6.27)$.

Fall 2: $T_U < T_F$ bei reiner Fixkostenbetrachtung

Bei einem höheren Steuersatz im Land U ($t_U > t_F$) $\rightarrow$ ($T_U < T_F$) unter reiner Fixkostenbetrachtung ($v = 0$) ist die Produktion von q_M unter CPM geringer als im Referenzszenario. (Detaillierter Beweis unter Anhang 1)

Ausgangsbedingung:

$$q_M (6.20) > q_M (6.27)$$

$$\frac{2}{7} \left(\frac{B(4 * K - 1)}{4} \right) > \frac{-(T_F - T_U) \frac{B}{4} + T_U (B \left(K - \frac{1}{4} \right))}{-(T_F - T_U) \frac{1}{4} + T_U \frac{7}{2}}$$

Diese Bedingung ist erfüllt, wenn $T_U < T_F$ ist.

Korrespondierend dazu steigt die Produktion (q_L) der lokalen Firma. Allerdings sinkt die Gesamtproduktionsmenge ($q_L + q_M$) gegenüber den Mengen im Referenzszenario, da wie in (6.26) ersichtlich, pro 1 Stück vermindertem q_M der Wert von q_L um 0,25 Stück steigt.

Zusätzlich erkennt man durch partielle Differenzierung der Differenzmenge $\Delta q_M(6.30) = q_M(6.20) - q_M(6.29)$ nach K , was ein negatives Ergebnis liefert, dass mit steigendem K der Wert von Δq_M sinkt, da mit jedem zusätzlichen K die Produktionsmenge von Division 2 unter CPM schneller wächst als im Referenzszenario.

$$\Delta q_M = \frac{2}{7} \left(\frac{B(4 * K - 1)}{4} \right) - \frac{-(T_F - T_U) \frac{B}{4} + T_U \left(B \left(K - \frac{1}{4} \right) \right)}{-(T_F - T_U) \frac{1}{4} + T_U \frac{7}{2}} \quad (6.30)$$

Fall 3: $T_U < T_F$ bei reiner Betrachtung der variablen Kosten

Auch bei reiner Betrachtung der variablen Kosten ($f = 0$) mit einem höheren Steuersatz im Land U ($t_U > t_F$) $\rightarrow (T_U < T_F)$ ist die Produktion von q_M unter CPM geringer als im Referenzszenario. (Detaillierter Beweis unter Anhang 2)

Ausgangsbedingung:

$$q_M(6.20) > q_M(6.29)$$

$$\frac{2}{7} \left(\frac{B(4 * K - 1)}{4} - \frac{3 * v}{4} \right) > \frac{(T_F - T_U) \frac{B - v}{2} + T_U \left(B \left(K - \frac{1}{4} \right) - \frac{3v}{4} \right)}{(T_F - T_U) + \frac{7}{2} T_U}$$

Diese Bedingung ist erfüllt, wenn $K > (2 - V/B)$ ist.

Korrespondierend dazu steigt die Produktion (q_L) der lokalen Firma L. Allerdings sinkt die Gesamtproduktionsmenge ($q_L + q_M$) gegenüber den CPM Mengen, da wie in (6.26) ersichtlich, pro 1 Stück vermindertem q_M der Wert von q_L um 0,25 Stück steigt.

Zusätzlich erkennt man durch partielle Differenzierung der Differenzmenge $\Delta q_M(6.31) = q_M(6.20) - q_M(6.29)$ nach K , was ein positives Ergebnis liefert, dass mit steigendem K der Wert von Δq_M steigt, da mit jedem zusätzlichen K die Produktionsmenge von Division 2 unter CPM langsamer wächst als im Referenzszenario.

$$\Delta q_M = \frac{2}{7} \left(\frac{B(4K - 1)}{4} - \frac{3v}{4} \right) - \frac{(T_F - T_U) \frac{B - v}{2} + T_U \left(B \left(K - \frac{1}{4} \right) - \frac{3v}{4} \right)}{(T_F - T_U) + \frac{7}{2} T_U} \quad (6.31)$$

Fall 4: $T_U > T_F$

Wenn der Steuersatz im Land U niedriger als im Land F ($t_U < t_F$) $\rightarrow (T_U > T_F)$ ist, ist der MNK grundsätzlich bestrebt, möglichst viel vom Profit in U anstatt in F zu versteuern. Um das zu erreichen, wird Division 2 unter CPM überproduzieren, wodurch die von L produzierte Menge q_L sowie deren Verkaufspreise abnehmen. Dadurch sinkt dann die Vergleichskennzahl Rohertrag geteilt durch betriebliche Aufwendungen des lokalen Unternehmens.

Es treten also genau die umgekehrten Effekte auf, wie die in den beiden zuvor beschriebenen Fällen.

Grundsätzliche Erkenntnisse des Szenarienvergleiches

Bei Einhaltung der in Fall 2 und Fall 3 von Abschnitt 6.2.3 aufgestellten Bedingungen steigt bei einem Rückgang von q_M die produzierte Menge des lokalen Unternehmens (q_L) leicht. Die Gesamtmenge ($q_M + q_L$) sinkt aber gegenüber der Gesamtmenge im Referenzszenario, was zu höheren Preisen für beide Unternehmen führt. Mit höherem Preis und erhöhter Produktionsmenge für das lokale Unternehmen im Vergleich zum Referenzszenario steigt natürlich auch der Profit für das lokale Unternehmen unter CPM. Das wiederum erhöht auch die Kennzahl Rohertrag durch betriebliche Aufwendungen für das lokale Unternehmen. Aus diesem Grund steigt auch der in F zu versteuernde Profit, was gleichzeitig den in U zu versteuernden Profit senkt und somit die Gesamtsteuerbelastung verringert..

In welcher Menge q_M unter CPM gegenüber dem Referenzszenario unterproduziert wird, wird durch die Stückzahl q_M bestimmt, ab der die Grenzprofitverringerung aufgrund der Profitverschiebung nach F bzw. der Nichtverschiebung nach U der Grenzsteuerersparnis entspricht.

Die genau umgekehrten Effekte treten auf, wenn ($t_U < t_F$) $\rightarrow (T_U > T_F)$ und dadurch Fall 4 aus Abschnitt 6.2.3 eintritt.

Falls es kein Abkommen gibt, das eine Doppelbesteuerung verhindert, wird der wirkliche Profit von Division 2 (Π_M) abzüglich des Werts von Π_{M2} (Die Differenz spiegelt den Wert des Immat. Wertgegenstandes wider.) in U und F versteuert. Die Profitfunktion des MNK (Π_Z) ändert sich auf $\{T_F * \Pi_M + T_U(\Pi_M - \Pi_{M2})\}$ unter Berücksichtigung derselben Bedingungen. Da so die Wertschöpfung des Immat. Wertgegenstandes doppelt besteuert wird, ist der

MNK natürlich umso mehr bestrebt, die Profitwirksamkeit des Immat. Wertgegenstandes durch Erhöhung des Profites von L zu senken, was durch die Unterproduktion von q_M erreicht wird. Dies ist auch der Fall, wenn $(t_U > t_F) \rightarrow (T_U < T_F)$ ist, weil es dann noch immer zu einer Doppelbesteuerung kommt.

6.2.4 Zusammenfassende Betrachtung der CPM-Analyse

Unabhängig davon, ob betriebliche Aufwendungen variabel, fix oder beides sind und ob das Land F ein Steuerabkommen mit U hat, welches die Doppelbesteuerung unterbindet oder nicht, hat die Verwendung von CPM folgenden Effekte:

- a. Bei Einhaltung der in Fall 2 und Fall 3 in Abschnitt 6.2.3 aufgestellten Bedingungen wird q_M im Vergleich zum Referenzszenario unterproduziert.
- b. Insgesamt betrachtet $(q_L + q_M)$ sinkt die produzierte und verkaufte Menge auf dem Markt im Land F.
- c. Die inhaltlichen Aussagen von a. und b. treffen weniger zu, umso größer K wird, weil damit Division 2 auf eine monopolistische Stellung zusteuert und somit das lokale Unternehmen an Bedeutung verliert.
- d. Die inhaltlichen Aussagen von a. und b. treffen mehr zu, wenn es kein Abkommen gibt, das eine Doppelbesteuerung verhindert. Dadurch wird das Steuereinsparungspotential vergrößert, welches durch die Minderung des TP für den Immat. Wertgegenstand entsteht.

6.3 Reflexion der gewonnenen Erkenntnisse

Im ersten Augenblick scheint die Vorgehensweise, die in den beiden behandelten Modellen angewendet wurde, sehr unproblematisch und einfach.

Allerdings muss dabei bedacht werden, dass der MNK die Herangehensweise der Behörden kennt und das zur Konzernprofitmaximierung nutzen kann. Sowohl bei der CPM- als auch beim CUT-Ansatz kann man erkennen, dass der MNK durch Beeinflussung der unkontrollierten Transaktion seinen eigenen Profit nach Steuern optimieren kann, wenn auch durch unterschiedliche Maßnahmen.

Inwieweit es korrekt ist, dass man wie bei der CPM Analyse die volle Profitdifferenz zwischen Division 2 und dem lokalen Unternehmen auf den Immat. Wertgegenstand zurückführt und somit die volle Profitdifferenz in U versteuert, darüber herrschen unterschiedliche Meinungen. Das amerikanische Gesetz gestattet dieses Vorgehen. DeSouza (1997, o.S.) führt aber an, dass man in so einem Fall nicht die volle Differenz auf die Marke (Immat. Wertgegenstand) zurückführen darf, da der innere Lizenznehmer (Division 2 in Abschnitt 6.2) die Risiken der kommerziellen Verwertung trägt und diese Risiken auch abgegolten werden müssen. Unterstützt wird diese Meinung durch mehrere amerikanische Gerichtsurteile (DeSouza, 1997, o.S.).

Darüber hinaus wird die Lizenzierung an ein unverbundenes Unternehmen nicht immer realisierbar sein bzw. meist nicht im Interesse des MNK liegen, was sich auf die in Abschnitt 3.2 beschriebenen Charakteristika zurückführen lässt. Wenn eine Lizenzierung aber an ein unverbundenes Unternehmen erfolgt, dann kann diese Transaktion zu einem inneren Fremdvergleich herangezogen werden, bei dem durch die Verwendung desselben Immat. Wertgegenstandes oft eine direkte Vergleichbarkeit besteht und so die notwendigen Anpassungen reduziert bzw. überhaupt unnötig werden.

Es ist aber wahrscheinlicher, wie in Abschnitt 4.8 beschrieben, dass sich aus mehreren unkontrollierten Vergleichstransaktionen mitunter auch mit verschiedenen Bestimmungsmethoden durch entsprechende Anpassungen eine Bandbreite (ALR) ermitteln lässt, in welcher der TP liegen muss. Bei Vorliegen einer solchen ALR wird der MNK den TP genauso wie unter CUT (Abschnitt 6.1) und CPM (Abschnitt 6.2) optimieren. Der MNK wird je nach Steuerdifferenz den TP auf den maximal zulässigen Wert aus der ermittelten Bandbreite (Grenze der ALR) festlegen. Er beeinflusst dadurch seine eigene unkontrollierte Transaktion nicht mehr, weil diese nicht mehr vorhanden ist bzw. zum Vergleich nicht herangezogen werden kann. Es tritt also der in Kapitel 5 von Boos (2003, S. 39ff) beschriebene Sachverhalt ein. Die von Halperin & Srinidhi (1996) gewonnenen Erkenntnisse haben also doch auch praktische Relevanz.

Diese praktische Relevanz wird durch eine Studie über amerikanische MNKs von Grubert & Mutti (2007, S. 1ff) bekräftigt, die aus den Ergebnissen ihrer Studie lesen, dass MNKs dazu tendieren, Divisionen in Niedrigsteuerländern

aufzubauen, um durch diese dann Patente und Lizenzen an produzierende Divisionen in Hochsteuerländern zu verkaufen.

Ähnliches gilt für MNKs in Europa. Dischinger & Riedel (2011, S. 691ff) konnten in ihrer Studie nachweisen, dass MNKs das Besitzrecht für Immat. Wertgegenstände zu Divisionen verschieben, die einem niedrigeren Steuersatz unterliegen.

Zusätzlich muss allerdings noch erwähnt werden, dass die CUT-Methode zur TP Bestimmung für Immat. Wertgegenstände in der Praxis wegen der fehlenden Vergleichstransaktionen wenig Relevanz besitzt (vgl. Brauner, 2008, S. 156).

Hinzu kommt, dass beim MNK nicht immer davon ausgegangen werden kann, dass das Ziel aller Entscheidungsträger die Profitmaximierung des Gesamtunternehmens ist. Wenn beispielsweise ein Manager in der Zentrale anhand des Wertbeitrages entlohnt wird, welcher durch die interne und externe Zurverfügungstellung des Immat. Wertgegenstandes gebildet wird, wird dieser Manager nicht danach trachten, den TP zu senken, um damit ebenfalls die Gesamtsteuerlast des MNK zu verringern. Dieser Sachverhalt wird aber in Kapitel 8 behandelt.

Außerdem kann dieses Problem, wie in Kapitel 7 beschrieben, durch die Einführung eines zusätzlichen, rein intern gültigen TP abgemildert werden, welcher zur Anreizsetzung verwendet wird.

7 TP FÜR IMMAT. WERTGEGENSTÄNDE BEI SQUENTIELLEM INVESTMENT

Dieses Kapitel beschäftigt sich mit Johnson (2006), die in ihrem Paper untersucht, welche Bedingungen für 3 unterschiedliche Methoden zur Bestimmung von TPs für immaterielle Wertgegenstände in einem bilateralen sequentiellen Investmentmodell gegeben sein müssen, um in Firmen mit dezentralisierter Organisationsstruktur korrekte Incentives für Investitions- und Handelsentscheidungen für Manager zu generieren.

Der Begriff Immat. Wertgegenstände umfasst eine ganze Reihe von Wertgegenständen mit zum Teil sehr unterschiedlichen Charakteristika. Johnson (2006, S. 340ff) trifft deshalb für ihr Modell eine Reihe von Annahmen. Zum Ersten kann der Wertgegenstand ex-ante, also vor der Erstellung nicht ausreichend genau beschrieben werden, um alle Details vor den Investitionen vertraglich festzulegen. Dieses Charakteristikum trifft auf eine Reihe von Immat. Wertgegenständen aus dem F&E Bereich zu. Als Beispiel kann hier die Entwicklung von Medikamenten oder Impfstoffen angeführt werden.

Zweitens können die Entwicklungskosten für Wertgegenstände von Dritten nicht beobachtet werden, weshalb keine Verträge auf Kostenbasis möglich sind.

Drittens können keine Opportunitätskosten auftreten, da der Immat. Wertgegenstand die Charakteristika eines öffentlichen Gutes besitzt und somit unendlich oft gleichzeitig genutzt werden kann. Als Beispiel kann man hier ein neu entwickeltes technisches Verfahren anführen.

7.1 Grundmodell

Johnson (2006) nimmt in ihrem Modell einen multinationalen Konzern (MNK) an, welcher zwei Tochterfirmen (Division 1 und Division 2) besitzt. Diese beiden Divisionen des MNK arbeiten zusammen, um gemeinsam einen Immat. Wertgegenstand zu entwickeln bzw. zu erzeugen, welcher die bereits zuvor beschriebenen Charakteristika aufweist.

Division 2 hat dabei ihren steuerpflichtigen Unternehmenssitz in einem Niedrigsteuerland mit Steuersatz t . Division 1 hat ihren steuerpflichtigen Unternehmenssitz in einem Hochsteuerland, welches den Steuersatz $t + h$ aufweist. Es herrschen also Steuerdifferenzen im Ausmaß von h .

Zur Entwicklung des Immat. Wertgegenstandes investiert Division 1 zum Zeitpunkt 1 den Betrag I_1 . Zum Zeitpunkt 2 kann es je nach Modell zu Verhandlungen über den Aufteilungsfaktor und somit über den TP kommen.

Daran anschließend wird der Immat. Wertgegenstand von Division 1 zu Division 2 transferiert. Aufgrund des öffentlichen Gut-Charakters könnte Division 1 den Wertgegenstand aber noch zu fremden Unternehmen transferieren. Dies ist jedoch von der Konzernzentrale untersagt, weil dadurch die vollständige Kontrolle über den Wertgegenstand verloren ginge und dies in weiterer Folge zu Umsatzeinbußen führen würde (vgl. Johnson, 2006, S. 344f).

Nachdem Division 2 den Immat. Wertgegenstand erhalten hat, investiert sie zum Zeitpunkt 3 den Betrag I_2 und wandelt somit den Immat. Wertgegenstand in ein verkaufsfertiges Gut um, welches dann zum Zeitpunkt 4 verkauft wird und einen operativen Profit von $M(I_1, I_2)$ erzielt. Praktische Beispiele für die Investition von Division 2 wären, Bewerben des Gutes, Hinzufügen eines neuen Features oder die Anpassung an spezielle Kundewünsche. Die nachfolgende Abbildung 2 veranschaulicht das Vorgehen anhand eines Zeitstrahls.

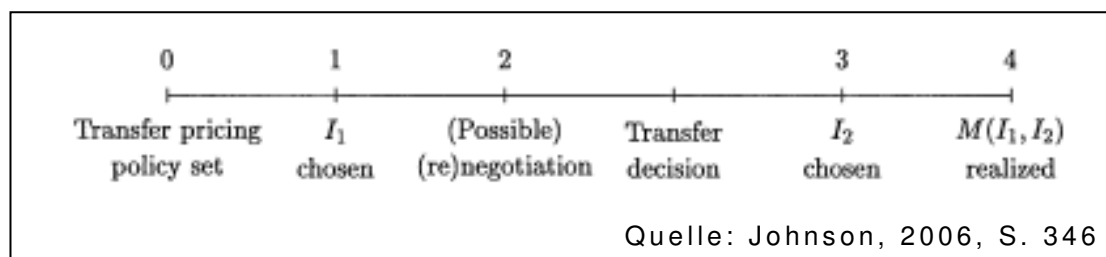


Abbildung 2: Modell-Zeitstrahl

Die TP-Bildung kann in unterschiedlicher Weise erfolgen. Die erste Art und Weise wäre die Festsetzung einer ex-ante Royalty Rate (auch als externer Aufteilungsfaktor bezeichnet) durch die Zentrale. Dadurch wird Division 2 veranlasst, den von der Zentrale definierten Prozentsatz des operativen Ergebnisses $M(I_1, I_2)$ an Division 1 zu überweisen. Der so entstandene TP für die konzerninterne Weitergabe des Immat. Wertgegenstandes ist mit der CPM bzw. TNM-Methode konform, da $M(I_1, I_2) = R(I_1, I_2, q) - v * q$ wobei $R(I_1, I_2, q)$ den Nettoumsatz, q die Stückzahl und v die Kosten pro q darstellen. Dabei wird q von Division 2 so gewählt, dass der operative Profit $M(I_1, I_2)$ maximiert wird.

7.1.1 Operativer Profit

Spezieller Beachtung bedarf noch die Funktion zur Generierung des operativen Profits $M(I_1, I_2)$. Da sie gleichzeitig von I_1 und I_2 abhängig ist, kann die Funktion 3 unterschiedliche Formen annehmen.

Zum Ersten können die Investments komplementär sein, sich also gegenseitig verstärken. Zum Zweiten können sich die Investments substituieren, sich also gegenseitig abschwächen. Und zum Dritten ist es möglich, dass sich I_1 und I_2 nicht gegenseitig beeinflussen. Dieses Verhalten bezeichnet Johnson (2006, S. 346f) also Quasi-Independence. Es stellt zugleich den einfachsten Fall für die spätere Betrachtung der Modelle dar, weil grundsätzlich dieselben Beobachtungen erzielt werden können, ohne dabei Effekte beachten zu müssen, welche durch die Wechselwirkung der Investitionen auftreten. (7.1) beschreibt diesen Quasi-Independence Zustand:

$$M(I_1, I_2) = \begin{cases} 0 & \text{if } I_1 < I_1^0 \text{ or } I_2 < I_2^0 \\ v(I_1) + w(I_2) & \text{if } I_1 \geq I_1^0 \text{ or } I_2 \geq I_2^0 \end{cases} \quad (7.1)$$

Das operative Ergebnis hängt von den Funktionen $v(I_1)$ und $w(I_2)$ ab, wobei beide Divisionen, wie in (7.1) ersichtlich ist, eine Mindestinvestition tätigen müssen. Grundsätzlich gilt, dass $v(I_1)$ und $w(I_2)$ positiv, fortwährend differenzierbar, stetig steigend und strikt konkav sein müssen, was auch für komplementäre und substituierende $M(I_1, I_2)$ gilt.

Im Endeffekt kann mit Quasi-Independence jede Division ihre Gewinne optimieren, ohne von der anderen Division abhängig zu sein (entsprechende Gewinnfunktion nach I_1 bzw. I_2 optimieren).

7.1.2 Transferpreis-Bestimmung

Der TP errechnet sich dann aus der Multiplikation des durch eine unkontrollierte Transaktion vorgegebenen Aufteilungsfaktors $\hat{a}$ ($0 \leq \hat{a} \leq 1$) mit dem $M(I_1, I_2)$, wie in Formel (7.2) ersichtlich ist.

$$TP = \hat{a} * M(I_1, I_2) \quad (7.2)$$

7.2 First Best

Betrachtet man dieses Modell aus der Perspektive der Zentrale des MNK, so spielt das sequentielle Investment der beiden Divisionen keine Rolle mehr, da die Zentrale unter Berücksichtigung der gegebenen Rahmenbedingungen das Investitionsniveau im Vorhinein schon optimal bestimmen kann. Dazu dient die folgende Profitfunktion (7.3):

$$\begin{aligned}\Pi(I_1, I_2 | h) = & (1 - t - h)[\hat{a} * M(I_1, I_2) - I_1] \\ & + (1 - t)[(1 - \hat{a}) * M(I_1, I_2) - I_2]\end{aligned}\quad (7.3)$$

Die Profitfunktion setzt sich dabei aus den Profitfunktionen von Division 1 (erste Zeile von (7.3)) und von Division 2 (zweite Zeile von (7.3)) zusammen. Gut zu erkennen ist dabei der TP $(\hat{a} * M(I_1, I_2))$, der dem zuerkannten operativen Profit von Division 1 entspricht. Davon abgezogen werden die Investitionsausgaben, wodurch der Gewinn vor Steuern entsteht. Dieser wird dann mit $(1 - t - h)$ multipliziert und ergibt somit den Profit nach Steuern von Division 1. Der Profit nach Steuern für Division 2 errechnet sich in gleicher Weise, nur muss Division 2 die Steuern in der Höhe von t entrichten und behält den Anteil $(1 - \hat{a})$ von $M(I_1, I_2)$.

7.3 Royaler Transferpreis bei sequentiellem Investment

7.3.1 Uniformer Royaler Transferpreis

In diesem ersten von drei Second Best-Modellen maximieren die beiden Divisionen sequentiell ihren Divisionsprofit anhand (7.4) und (7.5). Zum Zeitpunkt 2 kommt es zu keinen Verhandlungen, da der externe Aufteilungsfaktor $\hat{a}$ vor Beginn der Investitionen durch einen Fremdvergleich vorgegeben wird.

$$\Pi_1(I_1) = (1 - t - h)[\hat{a} (v(I_1) + w(I_2)) - I_1] \quad (7.4)$$

$$\Pi_2(I_2) = (1 - t)[(1 - \hat{a}) (v(I_1) + w(I_2)) - I_2] \quad (7.5)$$

Diese Formeln entsprechen zusammengesetzt der FB-Gesamtunternehmensprofitformel, wobei $M(I_1, I_2)$ jedoch in $v(I_1)$ und $w(I_2)$ aufgespalten wurde.

Das Ergebnis aus der FB-Betrachtung wird dadurch allerdings nicht erreicht. Klar ersichtlich wird das, wenn man die Profitfunktion der beiden Divisionen ansieht. Denn welche Investitionen eine einzelne Division getätigt hat, um den Wert von $M(I_1, I_2)$ zu steigern, wird dabei nicht berücksichtigt. Eine Division erhält also auch ihren prozentualen Anteil an $M(I_1, I_2)$, selbst wenn sie keine bzw. nur die Mindestinvestition tätigt. Die Divisionen erhalten nicht den ganzen Ertrag, den ihr Investment generiert und deshalb haben sie auch keinen Anreiz, bei einem Royalen TP mehr als nötig zu investieren.

Abbildung 3 soll diesen Sachverhalt besser veranschaulichen. Die X-Achse zeigt hier die Investitionshöhe, während die Y-Achse den Profit zeigt. Die blaue Profitfunktion zeigt dabei einen Aufteilungsfaktor $\hat{a}$ von 0,7 und die violette Profitfunktion zeigt einen Aufteilungsfaktor von 1. Die Division profitiert in diesem Fall also gänzlich von ihrer Investition. Klar ersichtlich ist in der Abbildung auch, dass dies zu einem höheren Investitionsniveau führt (maximales I bei $\hat{a} = 1$ größer als bei maximalem I bei $\hat{a} = 0,7$).

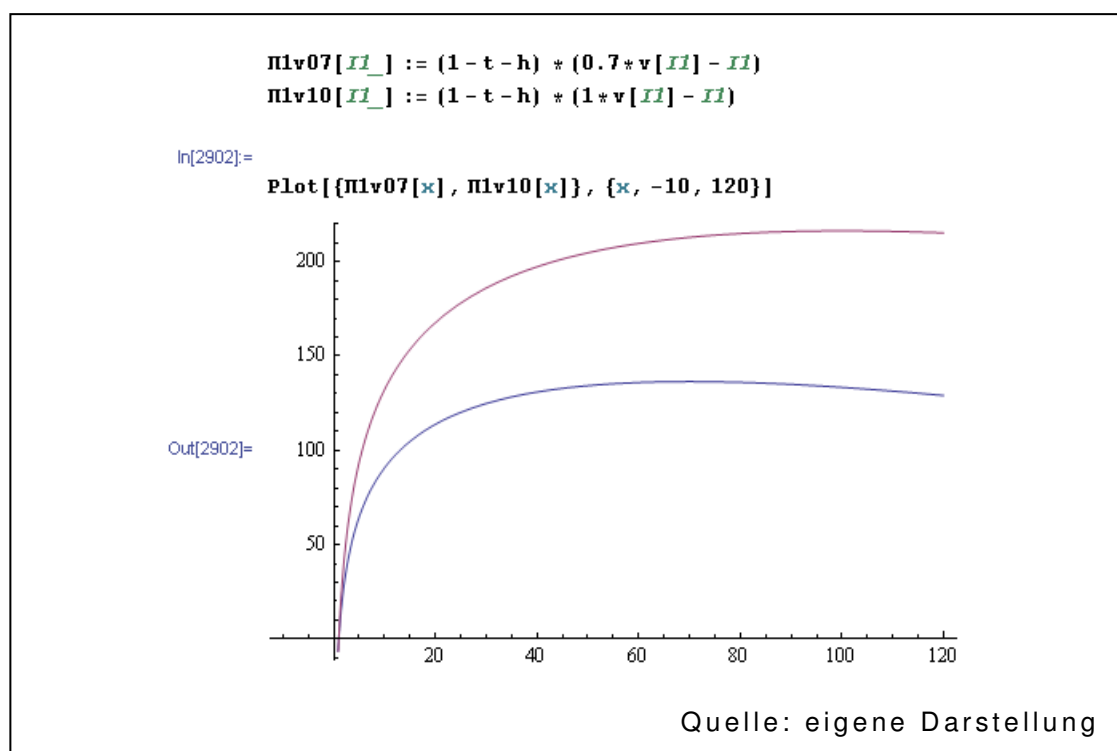


Abbildung 3: Investitionsniveaus bei unterschiedlichem Aufteilungsfaktor

Die Steuersätze selbst haben hier keinen Einfluss auf das Niveau der Investitionen, da sie erst auf den bereits maximierten Gewinn vor Steuern angewendet werden, welcher von der Investition herrührt.

Investitionsabweichung gegenüber FB

Hinsichtlich der Abweichung der Investitionen in der SB-Betrachtung im Vergleich zur FB-Betrachtung kann man erkennen, dass die Investitionsdifferenz von Division 1 zwischen FB und SB bei Quasi-Independence größer wird, da mit steigendem h Investition 1 (I_1) in FB größer wird und ein größeres h keine Anreize gibt, in einer SB-Situation die Investition zu steigern. Umso größer der Wert h in FB wird, umso mehr muss versteuert werden. Dieses Problem wird vermindert, indem man I_1 erhöht und somit den zu versteuernden Profit senkt.

Bei Division 2 tritt bei denselben Rahmenbedingungen genau der umgekehrte Effekt in Kraft. Bei steigendem h , sinkt I_2 in der FB-Situation und nähert sich der Investition an, welche Division 2 in der SB-Situation tätigt, da dort, wie bereits im Kapitel davor beschrieben, unterinvestiert wird.

Die Annahme, dass beide Divisionen wie bei Quasi-Independence unterinvestieren, gilt nicht für komplementäre und substituierende Investments, da dort je nach Rahmenbedingungen Division 1 zur Überinvestition neigt (vgl. Johnson, 2006, S. 349).

7.3.2 Decoupled Royaler Transferpreis

Neben der uniformen Transferpreissetzung, welche nach den gesetzlichen Rahmenbedingungen erfolgen muss und somit zu nicht optimalen Produktions- und Investment-Entscheidungen führen kann, gibt es auch noch die Möglichkeit, einen zusätzlich rein internen TP für die Steuerung und Anreizsetzung der einzelnen Divisionen zu verwenden. Dieses Verfahren bezeichnet man als Decoupling oder auch als „Keeping-two-sets-of-books“ (vgl. Johnson, 2006, S. 343)

Viele Firmen verzichten allerdings laut einer Studie von Ernst & Young (2003, S. 17) auf Decoupling, weil es höhere administrative Kosten generiert und sie befürchten, dass sie sonst stärker in den Focus der Steuerprüfer geraten.

Erweiterung des Modells mittels einem internen Aufteilungsfaktors

Das zuvor besprochene Modell kann auch für decoupele TP's verwendet werden. Der Profit von Division 1 berechnet sich dann wie in (7.6) und der Profit von Division 2 berechnet sich dann wie in (7.7) beschrieben.

$$\Pi_1^{\text{ro}} = a * M(I_1, I_2) - I_1 - (t + h)[\hat{a} * M(I_1, I_2) - I_1] \quad (7.6)$$

$$\Pi_2^{\text{ro}} = (1 - a)M(I_1, I_2) - I_2 - (t)[(1 - \hat{a})M(I_1, I_2) - I_2] \quad (7.7)$$

Neben dem offiziellen (externer) Aufteilungsfaktor $\hat{a}$, welcher den Profit nach gesetzlichen Regelungen zwischen den Divisionen aufteilt, wird jetzt noch zusätzlich der interne Aufteilungsfaktor a eingeführt ($0 \leq a \leq 1$), welcher in gleicherweise wie der externe Aufteilungsfaktor funktioniert. Wie in den vorherigen Formeln gut zu sehen ist, steht a in keinem Zusammenhang mit der Steuerermittlung.

Sehr wohl hängt der optimale interne Aufteilungsfaktor vom externen Aufteilungsfaktor $\hat{a}$ ab, da dieser konkrete Auswirkungen auf den Profit hat.

Es sei hier festgestellt, dass die Zentrale zur Maximierung des Profits den optimalen internen Aufteilungsfaktor bestimmen kann. Dieser hängt natürlich vom externen Aufteilungsfaktor $\hat{a}$ ab, da dieser konkrete Auswirkungen auf den Profit hat.

Um den Vorgang nicht unnötig zu verkomplizieren, wird Quasi-Independence mit $M(I_1, I_2) = v(I_1) + w(I_2)$ angenommen, wobei $v(I_1) = w(I_2)$ und $h = 0$ gilt. Durch die Symmetrie und $h = 0$ haben beide Divisionen die gleichen sinkenden Grenzerträge und werden deshalb in gleicher Höhe investieren ($I_1 = I_2$). Zur Berechnung des internen Aufteilungsfaktors a bildet man die FOC für die Profitfunktionen der beiden Divisionen ((7.6) und (7.7)) und setzt $v(I_1)$ und $w(I_2)$ gleich. Daraus lässt sich dann der optimale interne Aufteilungsfaktor berechnen. Unter Berücksichtigung der getroffenen Annahmen ergibt sich:

$$a^* = \frac{(1 - t)}{2} + \hat{a} * t \quad (7.8)$$

Der interne Aufteilungsfaktor a wird deshalb immer zwischen $\hat{a}$ und 0,5 liegen. Wenn $\hat{a}$ also bereits 0,5 ist, ergibt sich durch Decoupling kein Vorteil, weil auch

dann $\hat{a}$ den Wert 0,5 haben wird. Grundsätzlich kann man somit sagen, dass Decoupling den Profit nach Steuern hebt, je weiter $\hat{a}$ von 0,5 entfernt ist, weil durch den internen Aufteilungsfaktor insgesamt beide Divisionen in einem effizienteren Bereich der Profitfunktion operieren (fallender Grenzprofit). Dadurch ist es auch möglich, den Gesamtprofit zu erhöhen, selbst wenn keine Steuerunterschiede ($h = 0$) vorherrschen. Ebenso steigt natürlich der Gesamtprofit auch mit steigendem h .

Decoupling kurz zusammengefasst

Decoupling erhöht den Gesamtprofit des MNK, wenn $\hat{a}$ ungleich 0,5 ist. Das FB-Ergebnis wird dadurch aber nicht erreicht. Außerdem fallen zusätzliche Kosten durch erhöhten administrativen Aufwand an, welche aber in diesem Modell nicht berücksichtigt werden.

7.4 Verhandelter Transferpreis

Im Gegensatz zum Royalen TP-Modell bietet das Modell für einen Verhandelten TP die Möglichkeit, dass Division 1 und Division 2 zum Zeitpunkt 2 über die Höhe des internen Aufteilungsfaktors $\hat{a}$ verhandeln. Diese Verhandlungen finden beim Übergang des Immateriellen Wertgegenstandes von Division 1 zu Division 2 statt. Das bietet den Vorteil, dass konkrete Anreize für Division 2 geschaffen werden können, effektiv zu investieren. Als Nachteil steht dem gegenüber, dass Division 1 in eine Hold-Up-Problemsituation kommt, da sie bereits vor den Transferpreisverhandlungen ihre Investition tätigen muss.

Nachdem Division 1 ihre Investition getätigt hat, steht der folgende Profit aus (7.9) zur Verhandlung, wobei der Profit nach einer vorgegebenen Verteilungsregel verteilt wird. Im Endeffekt bekommt dann Division 1 den sogenannten γ -Anteil ($0 \leq \gamma \leq 1$). Dieser γ -Wert und somit die Verhandlungsstärke der Division kann zum Beispiel als Funktion der Verhandlungsfähigkeit der Divisionsmanager abgeleitet werden. Dies soll jedoch nicht Gegenstand dieser Arbeit sein. Respektive dazu verfügt der Manager von Division 2 die Verhandlungsstärke ($\gamma - 1$).

$$\max_{I_2} \{ [M(I_1, I_2) - I_2](1 - t) - [\hat{a} * h * M(I_1, I_2)] \} \quad (7.9)$$

Der zu verhandelnde Profit (7.9) ergibt sich aus dem kompletten operativen Profit, welcher durch die beiden Investitionen generiert wird. Abzüglich der

Investition von Division 2 ergibt das den verfügbaren Profit vor Steuern. Dieser wird dann noch mit Steuersatz t versteuert, welcher für beide Divisionen gilt. Zusätzlich müssen dann noch die Steuern, die nur für Division 1 anfallen, abgezogen werden.

Damit dieser Profit aber wirklich zur Verfügung steht, muss Division 2 effizient investieren. In bisherigen sequentiellen Modellen, wo bei den Verhandlungen der Besitz über das Verhandlungsobjekt von Division 1 an Division 2 übergeht, wird dies durch eine einfache Lumpsum Zahlung erreicht (vgl. Johnson, 2006, S. 351). In diesem Fall würde das aber falsche Anreize für Division 2 setzen, da die Steuerlast für Division 1, welche durch $M(I_1, I_2)$ generiert wird, nicht von Division 2 durch den TP entsprechend abgegolten wird.

Deshalb ist eine lineare Transferpreisregel notwendig, welche in (7.10) dargestellt wird:

$$TP = L + \beta * M(I_1, I_2) \quad (7.10)$$

β beschreibt dabei, welchen Anteil am Umsatz Division 1 zugewiesen bekommt. Zusätzlich wird dieser durch eine fixe Einmalzahlung L (Lumpsum) erhöht.

Dieser TP fungiert auch gleichzeitig als γ -Anteil von (7.9), wobei L vom Verhandlungsgeschick bestimmt wird und β vom Steuersatz und vom externen Aufteilungsfaktor $\hat{a}$, wie anschließend gezeigt wird. Daraus ergibt sich (7.11) für die Profitfunktion von Division 2 unter Verhandeltem TP.

$$\Pi_2^{\text{ver}} = M(I_1, I_2) - I_2 - L - \beta * M(I_1, I_2) - t[(1 - \hat{a}) * M(I_1, I_2) - I_2] \quad (7.11)$$

Damit die entsprechende Steuerlast mittels des TP an Division 1 weitergegeben wird, muss sich β aus der Steuerlast für Division 1, multipliziert mit dem externen Aufteilungsfaktor ergeben (siehe (7.12)). Es werden dadurch die Steuern, welche Division 1 resultierend aus ihrem Umsatzanteil bezahlen muss, über den TP von Division 2 an Division 1 weitergegeben.

$$\beta = \hat{a}(t + h) \quad (7.12)$$

Dieses Ergebnis für β in (7.12) kann formal noch hergeleitet werden, indem man vom Gesamtprofit bei zentraler Profitbestimmung (7.2) die FOC

hinsichtlich I_2 bildet und diese dann auf I_2 umformt und mit I_2 gleichsetzt. I_2 bekommt man, wenn man vom Profit von Division 2 unter Verhandeltem TP (7.11) die FOC nach I_2 bildet und auf I_2 umformt. Der Grund dafür liegt darin, dass I_2 in beiden Profitfunktionen hinsichtlich des gesamten MNK-Profites effizient ist.

Würde das nicht zutreffen und $\beta < \hat{a}(t + h)$ sein, so würde Division 2 aus Sicht der Zentrale überinvestieren, da Division 1 Steuerlasten übernehmen würde, die durch die Überinvestition von Division 2 entstanden sind und nicht durch den TP abgegolten werden. Division 2 könnte also so auf Kosten von Division 1 ihren Abteilungsprofit erhöhen.

7.4.1 Verhandelter TP kurz zusammengefasst

Bei einem Verhandeltem TP finden die Verhandlungen über den TP erst nach der Investition von Division 1 statt, wodurch Division 2 die Möglichkeit hat, sich Teile des Profits von Division 1 anzueignen. Antizipiert Division 1 dieses Vorgehen von Division 2, wird sie natürlich nicht effizient investieren. FB kann dadurch nicht erreicht werden, es sei denn, dass Division 1 schon vor der Tätigkeit von I_1 weiß, wie die Verhandlungen ausgehen werden, beispielsweise durch einen Vorabvertrag.

7.4.2 Vergleich von Verhandeltem und Royalem TP

Grundsätzlich kann man sagen, dass Division 2 bei Verhandeltem TP Vorteile gegenüber dem Royalem TP hat, da sie bei diesem TP immer effizient investieren wird. Bei Royalem TP dagegen wird Division 2 immer unterinvestieren.

Für Division 1 ist die ganze Sache schon etwas schwieriger, da diese bei einem Verhandeltem TP immer unterinvestieren und bei einem Royalem TP, je nach Verhalten der Investitionen zueinander, unter- oder überinvestieren wird. Dieser Fall tritt aber nicht ein, wenn die Investitionen von Division 1 und Division 2 voneinander unabhängig sind (Quasi Independence vorherrscht), da in diesem Fall Division 1 auch immer unterinvestieren wird.

Auf diese Art und Weise kann auch bestimmt werden, unter welchen Bedingungen der Verhandelte TP dem Royalem TP überlegen ist. Sobald Division 1 bei einem Verhandeltem TP mehr investiert als bei einem Royalem TP, ist der Verhandelte TP für die Zentrale besser als der Royale TP, da beim

Verhandelten TP Division 2 immer effizient investieren wird. Dabei wird hier von einem uniformen Royalen TP, also ohne Decoupling ausgegangen, da hier das Investitionsniveau von Division 1 höher ist als bei einem decoupelten Royalen TP. Vergleicht man das Investitionsniveau von Division 1 beim Verhandelten TP und dem Royalen TP, so entsteht folgende Bedingung (7.13). (Herleitung siehe Anhang 3). Wird diese erfüllt, ist der Verhandelte TP vorteilhafter für den MNK.

$$\frac{\gamma}{1 - (1 - \gamma) \frac{h}{1 - t}} > \hat{a} \quad (7.13)$$

Die Bedingung wird umso eher erfüllt, umso größer das Verhandlungsgeschick von Division 1 (γ) ist, da somit das Hold-up-Problem abgeschwächt wird. Je größer γ ist, umso mehr Profit wird Division 1 aus den Verhandlungen ziehen können. Zusätzlich steigt die Wahrscheinlichkeit, dass die Bedingung hält je kleiner der externe Aufteilungsfaktor $\hat{a}$ ist. Denn umso kleiner $\hat{a}$ ist, umso geringer ist der Anreiz für Division 1 bei einem Royalen TP zu investieren. Zusammengefasst heißt das, dass Division 1 die nötige Verhandlungsstärke braucht, um sich bei der Verhandlung die Umsätze zu sichern, welche ihr bei einem Royalen TP durch externen Aufteilungsfaktor $\hat{a}$ zustehen. Zusätzlich steigt die Wahrscheinlichkeit, dass die Bedingung hält, je größer die Steuendifferenz h ist, da mit steigendem h auch das Investitionslevel von Division 1 bei einem Verhandelten TP steigt. Im Gegenzug hat h keinen Einfluss auf das Investitionslevel bei einem uniformen Royalen TP.

7.5 Royaler TP mit Nachverhandlung

Mit der Verwendung des Royalen TP bei der Nachverhandlung zur Transferpreisbestimmung sollen die Schwächen der beiden von Johnson (2006) bisher vorgestellten Methoden behoben werden. Da es beim Royalen TP immer zur Unterinvestition von Division 2 kommt und beim Verhandelten TP immer ein Hold-up-Problem für Division 1 erzeugt wird, was zur Unterinvestition von Division 1 führt, sollen durch die Zusammenführung der beiden Systeme die gegenseitigen Schwächen ausgemerzt werden.

Die Zusammenführung findet in der Hinsicht statt, dass die beiden Divisionen ex ante einem Vertrag zustimmen, der aussagt, dass sich Division 1 bei Verhandlungen nicht verschlechtern darf. Division 1 steht dadurch keinem Hold-up-Problem mehr gegenüber. Der Profit, den Division 1 durch ihre eigene

Investition I_1 generiert, ist damit garantiert bzw. geschützt. Division 1 erhält mindestens den TP, der durch die Festsetzung des Royalen TP generiert wird, denn sonst würde sie dem späteren Verhandlungsergebnis nicht zustimmen.

Im ersten Schritt legt die Zentrale den internen Aufteilungsfaktor a fest. Im nächsten Schritt bestimmt Division 1 ihre Investition I_1 . Da Division 1 mit Sicherheit den Profit bekommt, welchen das Alpha des Royalen TP generiert, fungiert dieser quasi als Status quo. Der dadurch in den Verhandlungen zur Verfügung stehende Profit S errechnet sich wie in (7.14) beschrieben.

$$S = \{[M(I_1, I_2^{\text{ef}}) - I_2^{\text{ef}}](1 - t) - [\hat{a} * h * M(I_1, I_2^{\text{ef}})] - [M(I_1, I_2^{\text{ro}}) - I_2^{\text{ro}}](1 - t) - [\hat{a} * h * M(I_1, I_2^{\text{ro}})]\} \quad (7.14)$$

Die erste Zeile von (7.14) bildet den maximal noch zu erzielenden Profit ab, analog zu (7.9) (bei Verhandeltem TP zur Verfügung stehender Profit nach I_1), bei gegebenen I_1 und effizienter Investition von Division 2 (I_2^{ef}). Division 2 investiert effizient, wenn sie ihren Profit bei Verhandeltem TP (7.11) unter Berücksichtigung des TP aus (7.10) maximiert, und zwar aus denselben Gründen wie beim Verhandeltem TP angeführt. (7.15) zeigt, wie sich die effiziente Investition I_2^{ef} berechnet.

$$I_2^{\text{ef}} \in \max_{I_2} \{ \Pi_2^{\text{ver}} = M(I_1, I_2) - I_2 - L - \beta * M(I_1, I_2) - t[(1 - \hat{a}) * M(I_1, I_2) - I_2] \} \quad (7.15)$$

Die zweite Zeile von (7.14) bildet den noch zu erzielenden Profit, wenn die Royale TP-Bildung beibehalten wird. I_2^{ro} wird deshalb durch die Maximierung des Profites bei Royalem TP von Division 2 gebildet, wie in (7.16) beschrieben.

$$I_2^{\text{ro}} \in \max_{I_2} \{ \Pi_2^{\text{ro}} = (1 - a)M(I_1, I_2) - I_2 - t[(1 - \hat{a})M(I_1, I_2) - I_2] \} \quad (7.16)$$

Der Divisionsprofit für Division 1 unter dem Royalen TP mit Nachverhandlung ergibt sich somit, wie in (7.17) beschrieben, aus dem Profit, welcher unter dem Royalen TP entstanden wäre und dem γ -Anteil des durch die Nachverhandlung zur Verfügung stehenden Profites ($\gamma * S$).

$$\Pi_1 = a * M(I_1, I_2^{\text{ro}}) - I_1 - (t + h)[\hat{a} M(I_1, I_2^{\text{ro}}) - I_1] + \gamma * S \quad (7.17)$$

Wie in (7.18) beschrieben, berechnet sich der Profit von Division 2 analog zu (7.17), weil natürlich auch Division 2 einen Profit bei Royaler TP- Ermittlung erzielt hätte.

$$\Pi_2 = (1 - a)M(I_1, I_2^{ro}) - I_2^{ro} - t[(1 - \hat{a})M(I_1, I_2^{ro}) - I_2^{ro}] + (1 - \gamma)S \quad (7.18)$$

Mit dieser Vorgehensweise kann die Unterinvestition von Division 2 beim Royalen TP und das Hold-Up-Problem beim Verhandelten TP bei Division 1 verhindert werden. Aus diesem Grund dominiert der Royale TP mit Nachverhandlung die beiden anderen Methoden zur TP-Bildung.

Unter speziellen Voraussetzungen lässt sich dadurch auch das FB-Niveau erreichen. Dazu muss der interne Aufteilungsfaktor a von der Zentrale so groß gewählt werden, dass Division 1 so viel investiert, damit I_1 das FB-Niveau erreicht.

7.6 Reflexion der gewonnenen Erkenntnisse

Die Grunderkenntnis aus Johnson (2006) ist, dass bei sequentiellen Investments ein Royaler TP mit Nachverhandlungen einem reinen Royalen TP sowie einem Verhandelten TP überlegen ist.

Dafür verwendet Johnson (2006, S. 346) einen durch eine unkontrollierte Transaktion eindeutig vorgegebenen externen Aufteilungsfaktor. Aus Abschnitt 4.7 wissen wir, dass dieser eigentlich als Bandbreite angegeben werden muss, denn bei zumindest 2 oder mehreren Vergleichstransaktionen erhält man eine TP-Spannweite (ALR). Von den aus Kapitel 5 gewonnenen Erkenntnissen wissen wir aber auch, dass die Zentrale zur Profitmaximierung den kleinstmöglichen TP und somit den kleinstmöglichen externen Aufteilungsfaktor wählt, wenn für Division 1 der Steuersatz $(t + h)$ höher ist als der Steuersatz für Division 2 (t) . Für Johnsons (2006) Modell kann das implizit angenommen werden bzw. hat dieser Wert selbst keine Relevanz, da sein Einfluss in der FB- und SB-Situation gleich ist.

Ebenso wenig haben Steuern an sich Einfluss auf das Modell, da das Ergebnis durch Decoupling auch bei gleichen Steuersätzen (siehe Abschnitt 7.3.2) sowie bei Absenz von Steuern verbessert werden kann (Ausnahme externer Aufteilungsfaktor $\hat{a} = 0,5$).

Die Anwendung von Decoupling ist nicht auf Immat. Wertgegenstände beschränkt, wie zum Beispiel Hyde & Choe (2005, S. 166ff) zeigen. Gleichwohl kann das komplette Modell an sich auch für Mat. Wertgegenstände verwendet werden. Es muss nur von der Zentrale ein Verkauf nach extern untersagt werden, wodurch dann keine Opportunitätskosten mehr anfallen.

Somit ist Johnsons (2006) mehr dem Bereich des sequentiellen Investment zuzuordnen. Die Charakteristika des Immat. Wertgegenstandes werden dabei zu Hilfe genommen, um Opportunitätskosten, Vorverträge und auf Kosten basierende Verträge auszuschließen.

Allerdings stellt sich dann die Frage, inwieweit das sinnvoll ist, weil in so einem Fall der finale Wert des Wertgegenstandes konkreter beziffert werden kann und somit Verträge vor den Investitionen aufgesetzt werden können, welche die entsprechenden Investitionshöhen und Gewinnverteilungen regeln.

8 DEZENTRALE PROFITMAXIMIERUNG BEI VERHANDELTEM TP

In diesem Kapitel wird das von Dawson & Miller (2010) beschriebene Modell untersucht. Ziel dieses Modells ist es, unter dezentraler Profitmaximierung mit Verhandeltem TP denselben Profit wie unter zentraler Profitmaximierung zu erzielen. Dies soll erreicht werden, indem man den Bonus des Divisionsmanagers nicht mittels des jeweiligen Divisionsprofites bestimmt, sondern auf Basis der Summe aller Divisionsprofite. Dadurch werden die Divisionsmanager dazu veranlasst, alle relevanten Informationen miteinander zu teilen, um so gemeinsam den optimalen TP zu bestimmen, welcher den maximalen Gesamtprofit ermöglicht.

Zur einfacheren Erklärung des von Dawson & Miller (2010) vorgestellten Modells wird das in Kapitel 6 bereits verwendete Modell aufgegriffen. Der MNK verfügt darin über zwei Divisionen (Tochterunternehmen), die in unterschiedlichen Ländern operieren.

Division 2 stellt dabei das Tochterunternehmen dar, welches im Land F steuerpflichtig t_F ($T_F = 1 - t_F$) ist und dort auf dem Markt ein Produkt im Ausmaß von q_M absetzt. Der am Markt erzielbare Preis für das Produkt ergibt sich dabei wieder aus der inversen Nachfragefunktion ($K * B - q_M$), wobei K ausreichend groß ist, sodass die Teilnahmebedingung nicht greift.

Division 1 befindet sich im Land U, ist dort steuerpflichtig t_U ($T_U = 1 - t_U$) und übernimmt die Kontrolle über den Immat. Wertgegenstand. Division 1 ist aber verpflichtet, diesen Immat. Wertgegenstand der Division 2 zur Verfügung zu stellen, und zwar um den TP r pro verkauftem Stück q_M . Der Wert von r wird allerdings nicht vom MNK oder durch einen inneren Fremdvergleich vorgegeben, wie in Abschnitt 6 beschrieben. Der Bereich, in welchem sich der TP befinden muss, wird nun durch die ALR vorgegeben (siehe dazu Abschnitt 4.7). Diese ALR ist durch die Analyse mehrerer unkontrollierter Vergleichstransaktionen gegeben, welche entsprechend angepasst wurden. Sie reicht von der unteren Grenze r_{UG} bis zur oberen Grenze r_{OG} . Um den gesetzlichen Vorschriften zu entsprechen, muss r in diesem Bereich liegen. Es gilt also $r_{UG} \leq r \leq r_{OG}$.

Die in den Divisionen anfallenden Kosten werden jeweils mittels v (variable Kosten pro q_M) und f (fixe Kosten) dargestellt.

Darüber hinaus gibt es ein Abkommen zwischen den beiden Ländern, welches eine Doppelbesteuerung verhindert.

8.1 Zentrale und dezentrale Profitmaximierung

Um die Eigenschaften von Dawsons & Millers (2010) Modell besser veranschaulichen zu können, wird zuerst besprochen, wie sich der maximale MNK-Profit mittels Entscheidung durch die Zentrale (FB – Situation) vom Profit unter dezentraler Entscheidungsfindung (SB – Situation) unterscheidet.

8.1.1 Zentrale Profitmaximierung - FB

Bei zentraler Profitmaximierung ergibt sich der Profit (8.1) aus der Addition der beiden Divisionsprofitfunktionen ((8.2) & (8.3)).

$$\Pi_Z = T_u[q_M(r - v_1) - f_1] + T_F[q_M(K * B - q_M) - q_M(v_2 + r) - f_2] \quad (8.1)$$

Wie Kapitel 6 schon gezeigt hat, schafft r jedoch die Möglichkeit, die Gesamtsteuerlast durch Verschiebung des Profits in das Land mit dem niedrigeren Steuersatz zu minimieren und dadurch den Gesamtprofit zu maximieren.

Nimmt man an, dass die Steuerdifferenz zwischen den beiden Ländern so groß ist, dass der optimale TP jeweils außerhalb der ALR liegt (abgeleitet aus den Ergebnissen von Kapitel 6), so maximiert der MNK seinen Profit durch die Wahl jenes Grenzwertes der ALR (r_{UG} bzw. r_{OG}), der dem optimalen TP am nächsten liegt.

Zur einfacheren Veranschaulichung geht der Autor im weiteren Verlauf von den folgenden 2 Fällen aus:

Fall 1: $T_U < T_F$

Für den Fall, dass $t_U > t_F \rightarrow (T_U < T_F)$ gilt, wählt die Zentrale den niedrigsten durch die ALR zugelassenen Wert r_{UG} , damit möglichst wenig Profit in U versteuert werden muss.

Fall 2: $T_U > T_F$

Im Gegensatz dazu, wenn $t_U < t_F \rightarrow (T_U > T_F)$ gilt, wählt die Zentrale den höchsten durch die ALR zugelassenen Wert r_{OG} , um möglichst viel vom Profit in U zu versteuern.

8.1.2 Dezentrale Profitmaximierung mit Verhandeltem TP-SB

Bei dezentraler Profitmaximierung mit Verhandeltem TP verhandeln die beiden Divisionen zuerst über die Höhe von r und maximieren anschließend ihre Divisionsprofitfunktion, indem Division 2 die optimale Menge von q_M bestimmt.

Für Division 1 gilt die Profitfunktion (8.2). Sie besteht rein aus den Einnahmen, die durch die Lizenzierung des Immat. Wertgegenstandes an Division 2 erzielt werden, abzüglich der Kosten.

$$\Pi_1^{DZ} = T_u[q_M(r - v_1) - f_1] \quad (8.2)$$

Die variablen Kosten v_1 werden dabei mit dem Wert null angenommen, da alle Kosten, wie z.B. die jährlichen Patentgebühren als fix anzusehen sind und nicht von der Anzahl der Lizenzierungen (q_M) beeinflusst werden.

Der Profit von Division 2 errechnet sich, wie in Formel (8.3) ersichtlich, durch den Verkaufspreis ($K * B - q_M$) abzüglich der variablen Kosten v_2 und dem TP r multipliziert mit q_M minus der Fixkosten.

$$\Pi_2^{DZ} = T_F[q_M(K * B - q_M) - q_M(v_2 + r) - f_2] \quad (8.3)$$

Die Höhe der Steuersätze (t_F und t_U) haben bei der dezentralen Profitmaximierung keinen direkten Einfluss, sondern fungieren aus Sicht der Divisionsmanager als reine Dimensionierungsvariablen der Managerboni, welche sich vom jeweiligen Divisionsprofit ableiten ((8.4) für Managerbonus Division 1 und (8.5) für Managerbonus Division 2).

$$B_1^{DZ} = fB(\Pi_1^{DZ}) \quad (8.4)$$

$$B_2^{DZ} = fB(\Pi_2^{DZ}) \quad (8.5)$$

Der TP r spielt für die Profitmaximierung der beiden Divisionen eine Schlüsselrolle, da er für Division 1 die einzige Einnahmequelle darstellt. Für Division 2 ist r hingegen ein weiterer Kostenfaktor, den es zu minimieren gilt. Division 1 würde also den höchsten Wert favorisieren, den die ALR (r_{OG}) erlaubt, Division 2 dagegen den untersten Wert (r_{UG}).

Unabhängig davon, auf welchen Wert von r sich die beiden Divisionen in ihren Verhandlungen einigen können, wird dieser Wert r im Verhältnis der Verhandlungsstärke der beiden Divisionen immer innerhalb der ALR liegen.

Somit kann bei dieser Vorgehensweise in keinem der beiden unter FB betrachteten Fälle (Abschnitt 8.1.1) das FB-Ergebnis erreicht werden, außer in dem seltenen Fall, wenn der Divisionsmanager aus dem Niedrigsteuerland den anderen Divisionsmanager im Hochsteuerland in den Verhandlungen vollständig dominiert.

Die Verhandlungsstärke des Managers von Division 1 wird dabei wie in Kapitel 7 durch den γ -Anteil ($0 \leq \gamma \leq 1$) repräsentiert. Respektive dazu verfügt der Manager von Division 2 die Verhandlungsstärke $(1 - \gamma)$. Diese Werte lassen sich aus dem Verhandlungsgeschick der beiden Manager zueinander ableiten, was aber nicht Gegenstand dieser Arbeit ist.

8.2 Managerbonus erhöht sich durch Kooperation

8.2.1 Dezentrale Profitmaximierung mit Divisionskooperation

In diesem Abschnitt soll nun gezeigt werden, wie Dawson & Miller (2010) bei dezentraler Profitmaximierung die Manager der beiden Divisionen zur Kooperation motivieren kann und dadurch derselbe Profit wie bei zentraler Profitmaximierung erreicht werden kann. Der Wert von r kann, wie im vorherigen Abschnitt gezeigt wurde, innerhalb der ALR frei bestimmt werden.

(8.6) formuliert den Kooperationsprofit und stellt somit den Vorteil dar, welchen die Kooperation gegenüber der dezentralen Profitmaximierung (d.h. der jeweilige Divisionsmanager trachtet nach der Maximierung des eigenen Divisionsprofits) erreichen kann.

$$\Pi_{MNK}^{KV} = [\Pi_1^K + \Pi_2^K] - [\Pi_1^{DZ} + \Pi_2^{DZ}] \quad (8.6)$$

Um (8.6) zu maximieren, müssen die beiden kooperativen Divisionsgewinne in Summe dem Profit bei zentraler Profitbestimmung entsprechen ($[\Pi_1^K + \Pi_2^K] = \Pi_Z$). Um dies zu erreichen, muss der für den MNK optimale TP r und die daraus resultierende optimale Stückzahl q_M gewählt werden. Diese Werte entsprechen denen unter zentraler Profitmaximierung.

Damit die optimalen Werte (zum maximalen Profit führend) bestimmt werden können, müssen beide Manager von der Vorteilhaftigkeit dieses Modells überzeugt und bereit sein, alle dafür notwendigen privaten Informationen mit dem anderen Manager zu teilen.

Der Profit, der zur Bonusberechnung für den jeweiligen Divisionsmanager herangezogen wird, wird durch (8.7) bzw. (8.8) beschrieben:

$$\Pi_1^{BB} = \Pi_1^{DZ} + \gamma * \Pi_{MNK}^{KV} \quad (8.7)$$

$$\Pi_2^{BB} = \Pi_2^{DZ} + (1 - \gamma) * \Pi_{MNK}^{KV} \quad (8.8)$$

Der Profit, welcher der jeweiligen Division schon unter dezentraler Profitmaximierung zur Bonusberechnung gedient hat, wird also um den mit der Verhandlungsstärke der jeweiligen Division multiplizierten Kooperationsprofit aufsummiert.

Die Boni für die beiden Manager selbst berechnen sich dann wie im vorherigen Abschnitt durch die Anwendung der zuvor errechneten Profite (8.7) bzw. (8.8) auf die Bonusfunktion (8.9) bzw. (8.10).

$$B_1^{KV} = fB(\Pi_1^{BB}) \quad (8.9)$$

$$B_2^{KV} = fB(\Pi_2^{BB}) \quad (8.10)$$

Durch dieses Vorgehen verhandeln die beiden Divisionen um möglichst viel Profit für die eigene Abteilung erst dann, wenn der maximale Profit für den MNK schon feststeht. Die beiden Divisionsmanager arbeiten also vor den TP-Verhandlungen zusammen, damit der aufzuteilende Profit in den späteren Verhandlungen möglichst groß ist.

8.2.2 Probleme, welche sich durch die Divisionskooperation ergeben

Aus Sicht der Zentrale ist dieses Modell nicht ganz unproblematisch. Die Zentrale bekommt jetzt pro Division zusätzlich zum realen Divisionsprofit einen fiktiven, geplanten Profit (Π_1^{DZ} bzw. Π_2^{DZ}), wobei dessen Wichtigkeit mit abnehmender Verhandlungsstärke des Divisionsmanagers steigt, da dieser Profit für die Bonusberechnung voll wirksam ist und nicht noch im Verhältnis der Verhandlungsstärke geteilt werden muss.

Diese Profite (Π_1^{DZ} & Π_2^{DZ}) unter dezentraler Profitmaximierung sind deshalb fiktiv, weil sie nie in Wirklichkeit zustande gekommen sind. Die beiden Manager haben ja zusammengearbeitet und so den Kooperationsprofit erwirtschaftet.

Es besteht also ein Anreiz für die Divisionsmanager, ihren fiktiven Profit überhöht bzw. für die eigene Division maximal darzustellen, um dadurch ihren Bonus auf Kosten des anderen Managers zu erhöhen. Wie im Verlauf dieses Kapitels schon beschrieben wurde, ist der entscheidende Faktor dafür r .

Betrachtung von Fall 1 aus Abschnitt 8.1.1

Betrachtet man jetzt Fall 1 aus Abschnitt 8.1.1 und nimmt man dabei an, dass die beiden Divisionsmanager ihre fiktiven dezentralen Profite optimieren, kann man die Schwachstelle von Millers und Dawsons (2010) Modell erkennen.

Wie unter FB-Betrachtung festgestellt, liegt der für den MNK insgesamt optimale Wert von r bei r_{UG} . Dieser Wert wird auch auf Grund der Kooperation der beiden Divisionsmanager für die Bildung des Kooperationsprofites der jeweiligen Division (Π_1^K bzw. Π_2^K) herangezogen.

Ebenfalls wird Division 2 diesen TP zur Berechnung ihres fiktiven dezentralen Divisionsprofites heranziehen. Den Wert für q_M übernimmt der Manager aus der Kooperationsprofitberechnung oder er ermittelt ihn durch die Optimierung seiner Profitfunktion selbst. Beides ergibt das gleiche Ergebnis.

Dagegen wird der Manager von Division 1 r_{OG} für den TP ansetzen, da dieser höchstmögliche ALR-Wert den dezentralen Divisionsprofit von Division 1 maximiert (siehe Abschnitt 8.1.2). Es werden also zur Berechnung der fiktiven Divisionsprofite unterschiedliche Werte für r verwendet.

Zusammengefasst ergibt sich daraus, dass der fiktive dezentrale Divisionsprofit von Division 2 dem Kooperationsprofit von Division 2 entspricht ($\Pi_2^K = \Pi_2^{DZ}$) und dass der fiktive dezentrale Divisionsprofit von Division 1 höher als der Kooperationsprofit von Division 1 ist ($\Pi_1^K > \Pi_1^{DZ}$), da $r_{OG} > r_{UG}$.

Daraus resultiert, dass der Kooperationsvorteilsprofit aus (8.6) negativ ist ($\Pi_{MNK}^{KV} < 0$), da:

$$[\Pi_1^K + \Pi_2^K] < [\Pi_1^{DZ} + \Pi_2^{DZ}]$$

Die Bonusberechnungsbasis (8.7) bzw. (8.8) erhöht sich dadurch nicht. Ganz im Gegenteil, sie verringert sich sogar noch. Die Kooperation der beiden Divisionen bringt somit durch das opportunistische Verhalten des Managers von Division 1 keinen Vorteil und wirft dadurch die Frage nach der Sinnhaftigkeit dieses Vorgehens bzw. dieses Modells auf.

Betrachtung von Fall 2 aus Abschnitt 8.1.1

Dieser Fall 2 wird analog zu Fall 1 ablaufen, nur wird sich hier der Manager von Division 2 opportunistisch verhalten.

Verbesserungsansatz

Die Zentrale könnte zwar dieses opportunistische Verhalten aufdecken, wenn sie über die privaten Informationen der beiden Manager verfügen würde. Würde die Zentrale aber über diese Informationen verfügen, wäre zum einen die ganze dezentrale Profitmaximierung unnötig und zum anderen stellt sich die Frage, wie die unterschiedlichen Werte für r vereinheitlicht werden sollen? Im Grunde müssten die beiden Divisionsmanager einen für beide Divisionen gültigen TP aushandeln, und damit wäre man eigentlich beim Modell mittels dezentraler Profitbestimmung (Abschnitt 8.1.2) angekommen (FB wird darin nicht erreicht).

8.3 Dezentrale Profitmaximierung durch Decoupling

In diesem Abschnitt möchte der Autor dieser Arbeit die aus den vorangegangenen Kapiteln gewonnenen Erkenntnisse zusammenfügen und aufbauend auf die in Abschnitt 8.1 erklärten Grundlagen zeigen, wie unter bestimmten Voraussetzungen FB bei dezentraler Profitmaximierung erreicht werden kann.

Wie von Boos (2003, 39ff) bereits beschrieben, fungiert der TP aus Sicht der Zentrale im gesetzlichen Rahmen der ALR zum Verschieben der Steuern aus dem Hochsteuerland ins Niedrigsteuerland.

Daraus ergibt sich, dass je nach dem in welche Richtung der TP zur Verschiebung genützt wird, entweder der höchste oder der niedrigste Wert der ALR gewählt wird.

Demzufolge macht es auf längere Sicht gesehen Sinn, den Immat. Wertgegenstand ins Niedrigsteuerland zu verlagern, da so überhaupt kein TP ins Hochsteuerland überwiesen werden muss. Im umgekehrten Fall müsste

immer ein TP (kleinster Wert der ALR) ins Hochsteuerland überwiesen und dort versteuert werden (vgl Grubert & Mutti, 2007, S. 1ff; Dischinger & Riedel, 2011, S. 691ff).

8.3.1 Beschreibung der Ausgangslage

Wendet man diese Erkenntnisse auf das Modell in Kapitel 8 an, erhält man $t_U < t_F \rightarrow (T_U > T_F)$, was dem Fall 2 in Abschnitt 8.1.1 entspricht. In diesem Fall ist der höchste Wert der ALR r_{OG} , der profitmaximierende Wert aus Sicht der Zentrale. Daraus ergibt sich folgende Profitfunktion (8.11)

$$\begin{aligned} \Pi_Z = T_U[q_M(r_{OG} - v_1) - f_1] \\ + T_F[q_M(K * B - q_M) - q_M(v_2 + r_{OG}) - f_2] \end{aligned} \quad (8.11)$$

Bildet man die FOC hinsichtlich q_M und formt sie auf q_M um, erhält man die optimale Menge q_M ((8.12)).

$$q_M = \frac{T_F(B * K - r_{OG} - v_2) + T_U(r_{OG} - v_1)}{2 * T_F} \quad (8.12)$$

8.3.2 Dezentrale Profitmaximierung (SB)

Damit bei dezentraler Profitmaximierung der gleiche Profit wie bei zentraler Profitmaximierung erzielt wird, muss in jedem Fall der höchste Wert der ALR r_{OG} gewählt werden. Dadurch ergeben sich die Divisionsprofitfunktionen (8.13) und (8.14).

$$\Pi_1^{DZ} = T_U[q_M(r_{OG} - v_1) - f_1] \quad (8.13)$$

$$\Pi_2^{DZ} = T_F[q_M(K * B - q_M) - q_M(v_2 + r_{OG}) - f_2] \quad (8.14)$$

Wie in Abschnitt 8.1.2 beschrieben, favorisiert Division 1 unter normaler dezentraler Profitmaximierung einen möglichst hohen TP und Division 2 einen möglichst niedrigen TP (r_{OG}), da dieser erheblichen Einfluss auf den Divisionsprofit und somit auf den Bonus hat.

Vorgegebener TP

Bei der folgenden Vorgehensweise hat der Manager von Division 2 aber kein Problem mit diesem höchstmöglichen TP(r_{OG}), weil dieser im Endeffekt keinen Einfluss auf seinen Bonus hat. Die Gründe dafür werden im weiteren Verlauf erklärt.

Division 2 verwendet also r_{OG} als TP und kann somit den maximalen Profit von Division 2 durch das Festsetzen der optimalen Stückzahl q_M (8.15) bestimmen (FOC hinsichtlich q_M bilden und auf q_M umformen).

$$q_M = \frac{B * K - r_{OG} - v_2}{2} \quad (8.15)$$

Daraus lässt sich im Anschluss auch der Divisionsprofit von Division 1 bestimmen, da zu dessen Berechnung noch q_M (8.15) gefehlt hat. Der Gesamtprofit ergibt sich dann aus der Addition der beiden Divisionsprofite.

Dieser Gesamtprofit unter dezentraler Profitmaximierung (SB) ist kleiner als der Profit bei zentraler Profitmaximierung (FB), da die verkauften q_M in SB geringer sind als die q_M in FB unter der Berücksichtigung, dass $T_F < T_U$ ist, weil der TP nicht wirklich die anfallenden Grenzkosten unter Berücksichtigung der Steuerdifferenz widerspiegelt.

$$q_M (8.12) > q_M (8.15)$$

$$\frac{T_F(B * K - r_{OG} - v_2) + T_U(r_{OG} - v_1)}{2 * T_F} > \frac{B * K - r_{OG} - v_2}{2}$$

Decoupler TP für Division 2

Dieses Problem kann durch die Bildung des ökonomisch korrekten TP gelöst werden. Dieser muss gewährleisten, dass für den Immat. Wertgegenstand der Grenzertrag den Grenzkosten unter Berücksichtigung der Steuern entspricht (siehe Kap. 5). Für dieses Modell ergibt das den folgenden TP (8.16), der mit r_{SBO} bezeichnet wird.

$$r_{SBO} = \frac{r_{OG}(T_F - T_U) + v_1 * T_U}{T_F} \quad (8.16)$$

$(v_1 * T_U)$ stellt dabei die eigentlichen Grenzkosten des Immat. Wertgegenstandes unter Berücksichtigung der Steuern dar. Diese Grenzkosten können aber nicht als TP verwendet werden, da wie schon in Abschnitt 8.3.1 erläutert, r_{OG} der optimale TP aus Sicht der Zentrale ist. Deshalb wird r_{OG} unter Berücksichtigung der Steuerdifferenz $(T_F - T_U)$ zum Transferieren der Grenzkosten zu Division 2 verwendet. Die daraus resultierende Steuerreduzierung senkt die Grenzkosten, weshalb $r_{OG}(T_F - T_U)$ auch negativ ist, da $t_U > t_F \rightarrow (T_U < T_F)$. T_F unter dem Bruchstrich rechnet dann noch die Steuern ein, welche auf r_{SBO} durch Division 2 (Land F) wirken, weshalb man T_F auch kürzen kann wenn man r_{SBO} in (8.17) einsetzt.

Exkurs: Würde man in diesem Modell keine Steuern berücksichtigen, wäre r_{SBO} gleich v_1 .

Da angenommen wird, dass der Manager immer hinsichtlich seines eigenen Bonus maximiert, muss r_{SBO} Einfluss auf den Bonus des Managers haben. Das bedeutet, dass der Bonus des Managers von Division 2 nicht mehr vom wirklichen Profit abhängen darf, sondern von einem Dummyprofit (DZ_DP) (8.17), welcher den ökonomisch korrekten TP, r_{SBO} , verwendet.

$$\Pi_2^{DZ_DP} = T_F [q_M(K * B - q_M) - q_M(v_2 + r_{SBO}) - f_2] \quad (8.17)$$

Zur Bonusbestimmung decouplet man also den TP vom gesetzlich vorgeschriebenen TP (Decoupling wurde in Abschnitt 7.3.2 erklärt).

Der Manager von Division 2 bestimmt dann durch Errechnung der optimalen q_M (8.18) (FOC von (8.17) hinsichtlich q_M bilden und auf q_M umformen) den maximalen Dummyprofit.

$$q_M = \frac{T_F(B * K - r_{OG} - v_2) + T_U(r_{OG} - v_1)}{2 * T_F} \quad (8.18)$$

Gleichzeitig maximiert der Manager von Division 2 so seinen Bonus, der analog zu den Abschnitten 8.1.2 und 8.2.1 aus dem dafür herangezogenen Profit abgeleitet wird. Das dafür errechnete optimale q_M (8.18) entspricht dabei dem

q_M (8.12) aus der FB, wodurch auch die Summe der Divisionsprofite nach Steuern von Division 2 (8.14) und Division 1 (8.13) dem Profit aus FB (8.11) entspricht. Der Profit von Division 2 (8.14) allein betrachtet, ist dabei nicht maximal. Den Manager von Division 2 betrifft das nicht, da sein Bonus durch den Dummyprofit berechnet wird und nicht vom Divisionsprofit abhängig ist. Von der Zentrale ist das sogar gewünscht, da es vorteilhafter ist, den Profit im Niedrigsteuerland U ($(t_U < t_F) \rightarrow (T_U > T_F)$) zu versteuern.

Damit das Modell funktioniert, muss Division 2 wissen, wie sich r_{SBO} berechnet. Für die Berechnung sind der Division 2 alle Werte bekannt, ausgenommen v_1 , die variablen Kosten des Immat. Wertgegenstands, welche nur Division 1 exakt kennt (private Information von Division 1).

Das stellt aber kein wirkliches Problem dar, denn wie in Abschnitt 3.2 behandelt, erzeugt der Mat. Wertgegenstand, der als Träger des Immat. Wertgegenstandes fungiert, nur marginale Kosten im Vergleich zu dem Wert, den er schafft (Fix- bzw. Versunkene-Kosten wie Entwicklungskosten dürfen dabei nicht mehr berücksichtigt werden). Beispiele dafür sind Pillen, die einen medizinischen Wirkstoff enthalten, CDs, welche Software speichern oder einfach nur das Recht, eine bestimmte Marke führen zu dürfen ($v_1 = 0$, da kein Mat. Wertgegenstand benötigt wird). v_1 und damit r_{SBO} ist damit deutlich kleiner als r_{OG} .

Wenn der Mat. Wertgegenstand aufgrund seiner geringen, gegen null gehenden Kosten nicht ohnehin schon vernachlässigt werden kann oder wenn es sich um einen standardisierten Mat. Wertgegenstand handelt, dessen variable Kosten durch Marktbeobachtungen leicht zu ermitteln sind, dann kann FB bei dezentraler Profitmaximierung erreicht werden, indem Division 1 dafür sorgt, dass der MNK den optimalen ALR Wert wählt, weil dieser Wert auch den Bonus des Managers von Division 1 maximiert. Den Manager von Division 2 kümmert das nicht, weil dadurch der Profit, von dem sein Bonus abhängt, nicht beeinträchtigt wird, da dieser die variablen Kosten (Grenzkosten, inklusive steuerlicher Anpassungen) des Immat. Wertgegenstandes beinhaltet. In Summe führt das zum Erreichen des FB-Profites unter SB.

Dabei ist zu beachten, dass die Boni der beiden Manager auf Basis nicht voneinander abhängiger Profite berechnet werden. Dadurch ziehen beide Divisionen sozusagen am selben Strang, was für den MNK vorteilhaft ist, wenn

v_1 vernachlässigbar ist bzw. vom MNK einfach geprüft wird, ob die Annahmen von Division 2 für v_1 plausibel sind.

Für den außergewöhnlichen Fall, bei dem das nicht zutrifft und der Mat. Wertgegenstand selbst großen Wert besitzt bzw. dessen variable Kosten (v_1) schwer zu bestimmen sind, sollte dieses Vorgehen nicht angewendet werden. Der Manager von Division 2 wird dann möglichst niedrige v_1 (senkt Grenzkosten und steigert dadurch Produktion) annehmen, wodurch er den für ihn relevanten Dummyprofit erhöhen kann. Dieses Vorgehen wäre auch im Sinne des Managers von Division 1, weil Division 2 im Vergleich zur FB-Situation q_M überproduzieren wird und dadurch mehr Gewinne durch den vorgegebenen TP (r_{OG}) zu Division 1 verschoben wird, was somit den Bonus von Manager 1 erhöht. Beeinträchtigt wird dabei der wirkliche Profit von Division 2, welcher aber zu keiner Bonusberechnung herangezogen wird, sondern nur für den MNK relevant ist. Der FB-Profit wird durch die Überproduktion dann natürlich nicht erreicht.

Außerdem funktioniert dieses Vorgehen auch nicht, wenn $t_U > t_F \rightarrow (T_U < T_F)$ gilt, also der Immat. Wertgegenstand im Hochsteuerland angesiedelt ist. In diesem Fall präferiert nämlich der Manager von Division 1 trotzdem einen möglichst hohen TP (r_{OG}). Zum Erreichen von FB wäre aber r_{UG} notwendig.

Anmerkung:

Grundsätzlich ist es auch möglich, dass v_1 nicht nur durch die variablen Kosten des Mat. Wertgegenstandes definiert wird, sondern dass auch der Immat. Wertgegenstand variable Kosten aufweist und diese in v_1 einfließen. Ein Beispiel dafür wäre eine Logistiksoftware, die für jede Kundengruppe eigens neu entwickelt bzw. zum Teil neu entwickelt wird. In so einem Fall müsste dann aber auch die ALR jeweils neu bestimmt werden, weil es sich dabei immer um einen neuen Immat. Wertgegenstand handelt. Die variablen Kosten des Immat. Wertgegenstandes sind damit eigentlich Fixkosten und somit wäre das vorherige Vorgehensmodell wieder problemlos anwendbar.

Die hier angewandte Vorgehensweise lässt sich auch auf rein Materielle Wertgegenstände anwenden. Allerdings werden deren variable Kosten schwieriger zu bestimmen sein, wodurch die Wahrscheinlichkeit steigt, dass der Manager die variablen Kosten unterschätzt und seinen Bonus auf diese Weise erhöht. Der Profit des MNK jedoch sinkt dadurch. Außerdem ist im Normalfall

die ALR für Mat. Wertgegenstände kleiner, da mehr Vergleichstransaktionen zur Verfügung stehen. Das verringert auch ihr Steuerreduzierungs potenzial.

9 SCHLUSSKAPITEL

9.1 Zusammenfassung

Nach der grundlegenden Betrachtung des Nutzens bzw. der Vorteile, die der TP mit sich bringt, und der Charakterisierung von Immat. Wertgegenständen vor allem in ökonomischer Hinsicht, wurde im folgenden Schritt das Basisprinzip (Arm's length principle) erläutert, auf welchem alle OECD- und amerikanischen Methoden zur TP-Bestimmung beruhen.

In einem weiteren Schritt wurden die einzelnen Methoden theoretisch und beispielhaft erklärt und hinsichtlich ihrer Anwendbarkeit auf Immat. Wertgegenstände untersucht. Dabei wurde ersichtlich, dass die notwendigen Vergleichstransaktionen für die traditionellen Methoden meist nur eingeschränkt zur Verfügung stehen, wenn kein innerer Fremdvergleich möglich ist. Eine bessere Einsatzmöglichkeit finden im Gegensatz dazu die Gewinnmethoden, da sie keine Vergleichstransaktionen benötigen, in welcher der Immat. Wertgegenstand enthalten sein muss. Die aus den Vergleichen gewonnenen TPs ergeben dann zusammen die ALR, in welcher sich der gewählte TP schlussendlich befinden muss.

Bei der Untersuchung der Modelle von Boos (2003, S. 39ff) wurde die Erkenntnis gewonnen, dass der TP von Immat. Wertgegenständen prinzipiell dieselben ökonomischen Eigenschaften aufweist wie der TP für Mat. Wertgegenstände. Aufgrund der Tatsache, dass der TP innerhalb der ALR frei gewählt werden kann, entsteht die Möglichkeit, möglichst wenig Gewinn in Hochsteuerländer bzw. möglichst viel in Niedrigsteuerländer zu verschieben.

In einem weiteren Schritt wurde dann in Kapitel 6 gezeigt, dass es für einen MNK aus zentraler Betrachtung bei einem inneren Fremdvergleich sogar von Vorteil sein kann, seinen Profit aus einer unkontrollierten Transaktion zu schmälern, um dadurch den TP, der sich von dieser unkontrollierten Transaktion ableitet, für ihn vorteilhaft zu gestalten.

Inwieweit der Profit aus der unkontrollierten Transaktion geschmälert werden muss, damit der MNK seinen maximalen Gesamtprofit erwirtschaften kann, hängt von der Steuerdifferenz und vom Ausmaß ab, in welchem sich die Verkaufszahlen (Stückzahlen der unkontrollierten und kontrollierten Transaktion) durch den TP verändern.

Solche Umstände werden aber nicht immer gegeben sein, wobei man dann wieder bei der von Boos (2003, S. 39ff) beschriebenen Situation wäre, dass der MNK den optimalen Wert aus der für die kontrollierte Transaktion gebildete ALR bestimmt.

Da sich der TP für Immat. Wertgegenstände auch auf den Profit eines Unternehmens (Division) auswirkt, spielt er bei der Höhe der profitbezogenen Managerboni eine nicht zu unterschätzende Rolle. Die Steuerunterschiede spielen dabei für die Manager keine Rolle, da sie die Steuern nicht beeinflussen können. Auf Grund der Tatsache, dass der TP an sich für die liefernde Division den Ertrag und für die erhaltende Division die Kosten darstellt, favorisieren beide Divisionen die gegenteiligen Grenzwerte der ALR.

Um in so einer Situation die Möglichkeit zu schaffen, den gleichen Profit wie bei zentraler Profitmaximierung erreichen zu können, muss man grundsätzlich den zur Profitbestimmung herangezogenen TP vom steuerlich korrekten bzw. optimalen TP entkoppeln. In den Kapiteln 7 und 8 werden Lösungswege zur Anreizsetzung beschrieben, die es unter speziellen Voraussetzungen erlauben, den gleichen Profit wie bei zentraler Profitmaximierung zu erreichen.

9.2 Fazit

Zusammengefasst ergibt das die Schlussfolgerung, dass sich Mat. Wertgegenstände von Immat. Wertgegenständen nicht so sehr unterscheiden, weshalb auch theoretisch dieselben OECD Modelle zur TP-Bestimmung für Mat. und Immat. Wertgegenstände verwendet werden können. Alle in dieser Arbeit behandelten Modelle kann man deshalb auch für Mat. Wertgegenstände verwenden.

Allerdings bieten Immat. Wertgegenstände durch ihre schwierigere Bewertbarkeit eine größere Bandbreite zur Transferpreisfestsetzung (ALR), welche in Kombination mit ihrer unbegrenzten Wiederverwendbarkeit (Nichtrivalität) und ihren hohen Skalenerträgen ein viel größeres Steuersenkungs- bzw. Profitverschiebungspotential ermöglichen als dies Mat. Wertgegenstände können.

Des Weiteren wurde in dieser Arbeit auch auf das Spannungsfeld zwischen multinationalen Konzernen und den Steuerinteressen der verschiedenen Länder eingegangen. Ein Land ist immer bestrebt, die Geschäftstätigkeit dort zu erfassen und zu besteuern, wo sie geleistet wird. Ein Konzern, der international

agiert, will seine Steuern dort bezahlen, wo es für ihn am günstigsten ist. Mit Hilfe des TP können Gewinne zwischen Konzernteilen verschoben werden, um Steuerlasten zu reduzieren. Bei diesem Vorgang gilt es jedoch auch, Vorschriften und gesetzliche Regelungen zu beachten, auf die der Autor jedoch nur am Rande eingegangen ist.

Durch die Charakteristika des Immat. Wertgegenstandes in Verbindung mit dem Wunsch bzw. Willen der multinationalen Konzerne, Steuern zu sparen bzw. zu vermeiden, überwiegen beim TP für Immat. Wertgegenstände ganz klar die steuerlichen Aspekte, was in dieser Arbeit weitreichend dargestellt wurde.

Gleichzeitig wurde damit auch das grundlegende Modell erklärt, das die in der Einleitung erwähnten MNKs verwenden, um durch Profitverschiebung Steuern zu sparen.

LITERATURVERZEICHNIS

- Bakker, A. & Obuoforibo, B. (2009). *Transfer Pricing and Customs Valuation: Two worlds to tax as one*. Amsterdam: IBFD.
- Bianchi, P. & Labory, S. (2004). *The Economic Importance of Intangible Assets*. Burlington: Ashgate Publishing.
- Boos, M. (2003). *International Transfer Pricing: The Valuation of Intangible Assets*. Den Haag: Kluwer Law International.
- Brauner, Y. (2008). Value in the Eye of the Beholder: The Valuation of Intangibles for Transfer Pricing Purposes. *Virginia Tax Review*, 28(79), 79-164.
- Curd, R., Fickling, S. & Minnear, M. (2007). The Implications of Recent New York Transfer Pricing Decisions for IP Lawyers. *NYSBA Bright Ideas*, 16(1), 1-6. Zugriff am 12. Dezember 2012 unter <http://www.crai.com/uploadedFiles/Publications/implications-of-recent-new-york-transfer-pricing-decisions-for-ip-lawyers.pdf>.
- Dawson, P. C. & Miller S. M. (2010). Optimal Negotiated Transfer Pricing in Theory and Practice: Tangible and Intangible Intra-Firm Transfers. Zugriff am 12. Oktober 2012 unter <http://ideas.repec.org/p/uct/uconnp/2009-06.html>.
- DeSouza, G. (1997) Royalty methods for intellectual property. *Business Economics*, 32(2), Zugriff am 12. Jänner 2013 unter <http://www.freepatentsonline.com/article/Business-Economics/19545602.html>.
- Dischinger, M. & Riedel, N. (2011). Corporate taxes and the location of intangible assets within multinational firms. *Journal of Public Economics*, 95, 691-707.
- Ernst & Young (2010). 2010 Global Transfer Pricing Survey. Zugriff am 27. Dezember 2012 unter <http://www.ey.com/GL/en/Services/Tax/International-Tax/2010-Global-Transfer-Pricing-Survey>.
- Ernst & Young (2009). Transfer pricing enforcement takes center stage in the post-crisis environment. Zugriff am 21. Jänner 2013 unter

[http://www.ey.com/Publication/vwLUAssets/Transfer_pricing_enforcement/\\$FILE/Transfer%20pricing%20enforcement.pdf](http://www.ey.com/Publication/vwLUAssets/Transfer_pricing_enforcement/$FILE/Transfer%20pricing%20enforcement.pdf).

Ernst & Young (2005). 2005-2006 Global Transfer Pricing Surveys. Zugriff am 21. Jänner 2013 unter http://www.google.at/url?sa=t&rct=j&q=&esrc=s&source=web&cd=8&cad=rja&ved=0CFsQFjAH&url=http%3A%2F%2Fresources.rybinski.eu%2Fresources%2FsendFile%3Adb7fcc1c-be59-11de-85be-001b24eff4d8%3A1&ei=ipocUYnCDM_HtAar-YDABQ&usg=AFQjCNEExKuLqFQpwnTTDtaBtIsOv26e9NA&sig2=v931E3cx_hbSYZYGH3eVqQ&bvm=bv.42452523,d.Yms.

Ernst & Young (2003). Transfer Pricing 2003 Global Survey. Zugriff am 12. Oktober 2012 unter [http://webapp01.ey.com.pl/EYP/WEB/eycom_download.nsf/resources/Transfer%20Pricing%20Survey%20Report_2003.pdf/\\$FILE/Transfer%20Pricing%20Survey%20Report_2003.pdf](http://webapp01.ey.com.pl/EYP/WEB/eycom_download.nsf/resources/Transfer%20Pricing%20Survey%20Report_2003.pdf/$FILE/Transfer%20Pricing%20Survey%20Report_2003.pdf).

Gabler Wirtschaftslexikon (2012). Definition Verrechnungspreis. Zugriff am 12. Oktober 2012 unter <http://wirtschaftslexikon.gabler.de/Archiv/55441/verrechnungspreis-v5.html>.

Gravelle, J. G. (2013). Tax Havens: International Tax Avoidance and Evasion. Zugriff am 14. Februar 2013 unter <http://www.fas.org/sgp/crs/misc/R40623.pdf>.

Gu, F. & Lev, B. (2001). Markets in Intangibles: Patent Licensing. Zugriff am 10. Februar 2013 unter http://papers.ssrn.com/sol3/papers.cfm?abstract_id=275948.

Halperin, R. & Srinidhi B. (1996). U.S. Income Tax Transfer Pricing Rules for Intangibles as Approximations of Arm's Length Pricing. *The Accounting Review*, 71(1), 61-80.

Hyde C. E. & Choe C. (2005). Keeping Two Sets of Books: The Relationship Between Tax And Incentive Transfer Prices. *Journal of Economics & Management Strategy*, 14(1), 165-186.

Johnson, N. B. (2006). Divisional performance measurement and transfer pricing for intangible assets. *Accounting Studies*, 11(2-3), 339-365.

- Lev, B. (2001). *Intangibles: Management, Measurement, and Reporting*. Washington DC: Brookings Institution Press 2001.
- Mace, B., Faiferlick, C. & Levy, M. (2007). Financial Services Companies Increase Preparedness As Transfer Pricing Challenges Escalate. *Bank Accounting & Finance*, 20 (5), 29-34.
- Macho, R. & Perneki, M (2011). Verrechnungspreise: Benchmarking mittels Datenbankstudien – Fluch oder Segen?. Zugriff am 12. Dezember 2012 unter <http://www.pwc.at/files/publications/verrechnungspreise/verrechnungspreise-benchmarking-mit-datenbankstudien-swi-2011-07.pdf>.
- Markham, M. A. (2005). *The transfer pricing of intangibles*. Den Haag: Kluwer Law International.
- Markham, M. A. (2004). Transfer pricing of intangible assets in the US, the OECD and Australia: Are profit-split methodologies the way forward?. *University of Western Sydney Law Review*, 8, 55-78.
- Mutti, J. H. & Grubert, H. (2007). The Effect of Taxes on Royalties and the Migration of Intangible Assets Abroad. *Virginia NBER Working Paper Series*, Zugriff am 12. Jänner 2013 unter <http://www.nber.org/papers/w13248>.
- OECD (2012). DISCUSSION DRAFT - REVISION OF THE SPECIAL CONSIDERATIONS FOR INTANGIBLES IN CHAPTER VI OF THE OECD TRANSFER PRICING GUIDELINES AND RELATED PROVISIONS. Zugriff am 12. Oktober 2012 unter <http://www.oecd.org/ctp/transferpricing/50526258.pdf>.
- OECDb (2010). OECD Transfer Pricing Guidelines for Multinational Enterprises and Tax Administrations 2010. Zugriff am 12. September 2012 unter <http://www.oecd-ilibrary.org/content/serial/20769717>.
- OECDb (2010b). OECD ARM'S LENGTH RANGE. Zugriff am 2. Jänner 2013 unter <http://www.oecd.org/tax/transferpricing/45765058.pdf>.
- PWC (2012). International Transfer Pricing 2012. Zugriff am 1. Februar 2013 unter http://download.pwc.com/ie/pubs/2012_international_transfer_pricing.pdf.

- Ruf, M. (2007). *Steuerwettbewerb in Europa: Theorie, Empire und die Definition von Effektivsteuersätzen*. Berlin-Heidelberg: Springer Verlag.
- Schmid S. (2012). Die neue Macht des Konsumenten. Zugriff am 3. Februar 2013 unter <http://www.tagesanzeiger.ch/wirtschaft/unternehmen-und-konjunktur/Die-neue-Macht-des-Konsumenten/story/30471444>.
- Schreiber, U. (2008). *Besteuerung der Unternehmen - Einführung in Steuerrecht und Steuerwirkung*. Berlin-Heidelberg: Springer Verlag.
- Smith, G. V. & Parr, R. L. (2005). *Intellectual Property: Valuation, Exploitation and Infringement Damages 4th Edition*. Hoboken: Wiley & Sons.
- Szigetvari A. (2013). OECD attackiert Steuer-Schlupflöcher. Zugriff am 13. Februar 2013 unter <http://derstandard.at/1360681316700/OECD-attackiert-Steuer-Schlupfloecher-der-Konzerne?seite=2#forumstart>.
- Szigetvari A. (2012). Mit Google über die Bermudas surfen. Zugriff am 3. Februar 2013 unter <http://derstandard.at/1355459809311/Mit-Google-ueber-die-Bermudas-surfen>.
- Urquidi, A. J. (2008). AN INTRODUCTION TO TRANSFER PRICING. *New School Economic Review*, 3(1), 27-45.
- Verhülsdonk (2009). Internationale Verrechnungspreise und ihre Dokumentation. Zugriff am 28. September 2012 unter http://www.verhuelsdonk.de/pub/info/downloads/bro_verrechnungspreise.pdf.
- VPR (2010). Verrechnungspreisrichtlinien 2010. Zugriff am 12. Dezember 2012 unter <https://findok.bmf.gv.at/findok/link?gz=%22BMF-010221%2F2522-IV%2F4%2F2010%22&gueltig=20101028&bereich=rl>.

ANHANG 1

$$q_M (6.20) > q_M (6.27)$$

$$\frac{2}{7} \left(\frac{B(4 * K - 1)}{4} - \frac{3 * v}{4} \right) > \frac{-(T_F - T_U) \frac{B}{4} + T_U (B \left(K - \frac{1}{4} \right))}{-(T_F - T_U) \frac{1}{4} + T_U \frac{7}{2}}$$

Da $v = 0$, vereinfacht sich der Vergleich auf:

$$\frac{2}{7} \left(\frac{B(4 * K - 1)}{4} \right) > \frac{-(T_F - T_U) \frac{B}{4} + T_U (B \left(K - \frac{1}{4} \right))}{-(T_F - T_U) \frac{1}{4} + T_U \frac{7}{2}}$$

$$\frac{2}{7} \left(B \left(K - \frac{1}{4} \right) \right) > \frac{-(T_F - T_U) \frac{B}{4} + T_U (B \left(K - \frac{1}{4} \right))}{-(T_F - T_U) \frac{1}{4} + T_U \frac{7}{2}}$$

Zur Vereinfachung werden die folgenden Terme als Variablen definiert:

$$a = \frac{(T_F - T_U)}{4} \quad b = T_U * B \left(K - \frac{1}{4} \right) \quad d = T_U \frac{7}{2}$$

Daraus ergibt sich:

$$q_M (6.20) > q_M (6.27)$$

$$\frac{b}{d} > \frac{b - B * a}{d - a}$$

$$\frac{b}{d} > \frac{b - B * a + B * d - B * d}{d - a}$$

$$\frac{b}{d} > \frac{B(d - a)}{d - a} + \frac{b \frac{d}{d} - B * d}{d - a}$$

$$\frac{b}{d} > B + \frac{d}{d - a} \left(\frac{b}{d} - B \right)$$

$$0 > \left(B - \frac{b}{d}\right) - \frac{d}{d-a} \left(B - \frac{b}{d}\right)$$

$$0 > \left(B - \frac{b}{d}\right) \left(1 - \frac{d}{d-a}\right)$$

Aus (6.19) erkennt man, dass bei $v = 0$ der Wert von q_M ((6.20) = b/d) kleiner als B sein muss, damit q_L nicht negativ werden kann. Somit ist der erste Klammerausdruck in der Bedingung oberhalb dieses Textes positiv. Aus diesem Grund muss zur Erfüllung dieser Bedingung der zweite Klammerausdruck negativ sein. Da d in jedem Fall positiv ist, muss a auch positiv sein. Da gilt, dass $T_F > T_U$ ist, ist auch a positiv und somit ist die Bedingung erfüllt.

ANHANG 2

$$q_M (6.20) > q_M (6.29)$$

$$\frac{B(4 * K - 1)}{14} - \frac{3 * v}{14} > \frac{(T_F - T_U) \frac{B - v}{2} + T_U \left(B \left(K - \frac{1}{4} \right) - \frac{3v}{4} \right)}{(T_F - T_U) + \frac{7}{2} T_U}$$

$$\frac{2}{7} \left(\frac{B(4 * K - 1)}{4} - \frac{3 * v}{4} \right) > \frac{(T_F - T_U) \frac{B - v}{2} + T_U \left(B \left(K - \frac{1}{4} \right) - \frac{3v}{4} \right)}{(T_F - T_U) + \frac{7}{2} T_U}$$

$$\frac{T_U \left(B \left(K - \frac{1}{4} \right) - \frac{3 * v}{4} \right)}{\frac{7}{2} T_U} > \frac{(T_F - T_U) \frac{B - v}{2} + T_U \left(B \left(K - \frac{1}{4} \right) - \frac{3v}{4} \right)}{(T_F - T_U) + \frac{7}{2} T_U}$$

Zur Vereinfachung werden die folgenden Terme als Variablen definiert:

$$d = T_U \frac{7}{2} \qquad e = (T_F - T_U)$$

$$f = T_U * \left[B \left(K - \frac{1}{4} \right) - v \frac{3}{4} \right] \qquad g = \frac{(T_F - T_U)(B - v)}{2}$$

Daraus folgt:

$$q_M (6.20) > q_M (6.29)$$

$$\frac{f}{d} > \frac{f + g}{d + e}$$

$$f * d + f * e > d * f + d * g$$

$$f * e > d * g$$

$$T_U * \left[B \left(K - \frac{1}{4} \right) - v \frac{3}{4} \right] * (T_F - T_U) > \frac{(T_F - T_U)(B - v)}{2} * T_U \frac{7}{2}$$

$$\left[B \left(K - \frac{1}{4} \right) - v \frac{3}{4} \right] > \frac{(B - v)}{2} * \frac{7}{2}$$

$$BK - \frac{B}{4} - v\frac{3}{4} > \frac{7 * B}{4} - \frac{7 * v}{4}$$

$$BK - \frac{8 * B}{4} > -\frac{4 * v}{4}$$

$$B(K - 2) > -v$$

$$K > 2 - \frac{v}{B}$$

Wenn diese Bedingung erfüllt wird, gilt: $q_M(6.20) > q_M(6.29)$

ANHANG 3

Royaler TP im Vergleich zum Verhandelten TP

Indem man die beiden TP-Modelle vergleicht, gelangt man zu einer hinreichenden Bedingung, die bei Erfüllung zeigt, dass der Verhandelte TP dem Royalen TP aus Sicht der Zentrale überlegen ist.

Für das Royale TP-Modell verwendet man die Profitfunktion von Division 1 unter einem uniformen Royalen TP, da hier das Investitionslevel von Division 1 höher als das Investitionslevel unter einem decoupelten Royalen TP ist ($\hat{a} > 0,5$). Außerdem muss dann kein optimaler interner Aufteilungsfaktor berechnet werden, wobei dieser, falls er dennoch berechnet wird, bei unterschiedlichen Funktionen für $v(I_1)$ natürlich auch anders aussehen würde. Deshalb kann dann auch kein allgemein gültiges Ergebnis erreicht werden.

Uniformer Royler TP - I_1

Man nimmt die Divisionsprofitfunktion von Division 1 unter uniformen Royalen TP,

$$\Pi_1^{\text{Roy}}(I_1, I_2) = (1 - t - h)[\hat{a} * M(I_1, I_2) - I_1]$$

bildet die FOC hinsichtlich I_1 und formt sie auf $v'(I_1)$ um,

$$v'^{\text{Roy}}(I_1) = \frac{1}{\hat{a}}$$

womit man das optimale I_1 für eine unspezifizierte Funktion $v(I_1)$ unter Royalem TP erhält ($v'^{\text{Roy}}(I_1)$).

Verhandelter TP - I_1

Beim Verhandelten TP maximiert Division 1 ihren Profit, indem sie vom γ -Anteil des in den Verhandlungen zur Verfügung stehenden Profites die Investition abzieht und das ganze korrekt versteuert. Man multipliziert also γ mit (7.9) und zieht davon die um den Steuerschild bereinigte Investition I_1 ab.

$$\Pi_1^{\text{Ver}}(I_1, I_2) = \gamma[M(I_1, I_2)(1 - t - \hat{a} * h) - I_2(1 - t)] - I_1(1 - t - h)$$

Davon bildet man die FOC hinsichtlich I_1 und formt sie auf $v'(I_1)$ um und erhält dadurch:

$$v'^{\text{Ver}}(I_1) = \frac{(1 - t - h)}{\gamma(1 - t - \hat{a} * h)}$$

Womit man das optimale I_1 für eine unspezifizierte Funktion $v(I_1)$ unter Verhandeltem TP erhält ($v'^{\text{Ver}}(I_1)$).

Formuliert man nun die Bedingung, sodass das Investitionsniveau beim Verhandeltem TP höher als beim Royalen TP liegt, erhält man:

$$v'^{\text{Ver}}(I_1) > v'^{\text{Roy}}(I_1)$$

Setzt man die oberhalb errechneten Ausdrücke ein, ergibt sich:

$$\frac{(1 - t - h)}{\gamma(1 - t - \hat{a} * h)} > \frac{1}{\hat{a}}$$

Zur einfacheren Interpretation formt man die Bedingung mit den folgenden Schritten um:

$$\frac{\hat{a}(1 - t - h)}{1} > \frac{\gamma(1 - t - \hat{a} * h)}{1}$$

$$\hat{a} - \hat{a} * t - \hat{a} * h > \gamma - \gamma * t - \gamma * \hat{a} * h$$

$$\gamma * t - \gamma > \hat{a} * t - \hat{a} + \hat{a} * h - \gamma * \hat{a} * h$$

$$\gamma(t - \gamma) > \hat{a}(t - 1 + h - \gamma * h)$$

$$\frac{-1}{-1} * \frac{\gamma(t - \gamma)}{(t - 1 + h - \gamma * h)} > \hat{a}$$

$$\frac{\gamma(1 - t)}{(1 - t - h + \gamma * h)} > \hat{a}$$

$$\frac{\gamma}{\frac{(1 - t - h + \gamma * h)}{(1 - t)}} > \hat{a}$$

$$\frac{\gamma}{1 - \frac{(h - \gamma * h)}{(1 - t)}} > \hat{a}$$

$$\frac{\gamma}{1 - (1 - \gamma) \frac{h}{(1 - t)}} > \hat{a}$$

Dieses Ergebnis entspricht (7.13)

ABSTRAKT

Diese Arbeit befasst sich mit Transferpreisen für Immaterielle Wertgegenstände und betrachtet dabei sämtliche Publikationen, welche sich mit ökonomischen Transferpreismodellen in Bezug auf immaterielle Wertgegenstände beschäftigen.

Als Grundlage dafür werden die Eigenschaften und Charakteristika von Transferpreisen und Immateriellen Wertgegenständen erläutert. Darauf beruhend werden die gesetzlichen Rahmenbedingungen, welche für die Transferpreisbestimmung von Immateriellen Wertgegenständen gelten, in theoretischer und praktischer Weise erklärt.

Darauf aufbauend werden dann die einzelnen Modelle untersucht, mit dem Ergebnis, dass Transferpreise für Immaterielle Wertgegenstände auf der gleichen Ökonomie basieren wie Transferpreise für Materielle Wertgegenstände.

Allerdings können Immaterielle Wertgegenstände aufgrund ihrer Charakteristika viel schwieriger bewertet werden und bieten deshalb größeres Potenzial zur Steuervermeidung, weil multinationale Konzerne ihre Steuern dort bezahlen wollen, wo es für sie am günstigen ist. Mithilfe von Transferpreiszahlungen für Immaterielle Wertgegenstände können sie effektiv Gewinne zwischen ihren Konzernteilen verschieben und dadurch ihre Gesamtsteuerlast reduzieren.

Des Weiteren wurde gezeigt, dass durch die Einführung eines zweiten TP, welcher vom steuerlich korrekten bzw. optimalen TP entkoppelt ist und so für die interne Anreizsetzung optimiert werden kann, d.h. bei dezentraler Profitmaximierung der gleiche Profit wie bei zentraler Profitmaximierung erreicht werden kann.

LEBENS LAUF

Name: Christoph Schwertl

Nationalität: Österreich

E-Mail: a1007247@unet.univie.ac.at

Ausbildung:

- | | |
|-------------------|--|
| 2010.09 – 2013.06 | Universität Wien, Masterstudium: Betriebswirtschaft -
Spezialisierungen: Controlling, Wirtschaftsinformatik. |
| 2007.09 – 2010.06 | IMC FH Krems, Bachelorstudium:
Unternehmensführung & E-Business Management
Spezialisierungen: Gründungsfinanzierung. |
| 2000.09 – 2005.06 | HTBLA Salzburg für Elektronik: Spezialisierung:
Technische Informatik, Abschluss mit Reife- und
Diplomprüfung. |