



URZĄD OCHRONY KONKURENCJI I KONSUMENTÓW

Instrumenty ekonomiczne w prawie konkurencji



Publikacja przygotowana
dzięki wsparciu finansowemu
Unii Europejskiej

Instrumenty ekonomiczne w prawie konkurencji



URZĄD OCHRONY KONKURENCJI I KONSUMENTÓW

Instrumenty ekonomiczne w prawie konkurencji

Christian Bongard
Dirk Möller
Achim Raimann
Nikodem Szadkowski
Urszula Dubejko



Bonn / Warszawa 2007

Niniejsza publikacja jest współfinansowana ze środków Unii Europejskiej w ramach projektu Transition Facility 2004 (PL/IB/FI/02) „Ochrona Konkurencji”

ISBN 978-83-60632-05-5

Urząd Ochrony Konkurencji i Konsumentów
Plac Powstańców Warszawy 1
00-950 Warszawa
tel. 22 55 60 800
www.uokik.gov.pl

Warszawa 2007

Spis treści

1	Wprowadzenie	5
1.1	Jak korzystać z tego podręcznika	5
1.2	Zastosowanie metod analizy ekonomicznej w praktyce prawa konkurencji ..	6
2	Analiza mikroekonomiczna	7
2.1	Podstawy	7
2.2	Modele	10
2.2.1	Założenia ogólne	10
2.2.2	Analiza regresji	11
3	Rynek	15
3.1	Strona popytu.....	15
3.1.1	Popyt indywidualny i popyt rynkowy	15
3.1.2	Elastyczność popytu	18
3.1.3	Nadwyżka konsumenta.....	19
3.2	Strona podaży.....	20
3.2.1	Podaż indywidualna i podaż rynkowa	20
3.2.2	Koszty i zyski przedsiębiorstwa	22
3.2.3	Nadwyżka producenta	25
3.3	Zależność między podażą a popytem.....	26
3.4	Fazy rynku	28
4	Formy rynku i kształtowanie cen.....	30
4.1	Założenia ogólne.....	30
4.2	Konkurencja doskonała (polipol)	32
4.3	Monopol	33

4.4	Konkurencja monopolistyczna	35
4.5	Rynki oligopolistyczne	36
4.5.1	Ogólnie	36
4.5.2	Teoria gier	37
4.5.3	Model Cournota (konkurencja ilościowa)	40
4.5.4	Model Bertranda (konkurencja cenowa)	43
4.5.5	Rynki niejednorodne (heterogeniczne)	46
5	Zawodność rynku.....	48
5.1	Efekty zewnętrzne	49
5.2	Niepodzielność	50
5.2.1	Założenia ogólne	50
5.2.2	Korzyści skali (economies of scale)	51
5.2.3	Korzyści zakresu (economies of scope)	52
5.2.4	Skończoność korzyści skali i zakresu	53
5.2.5	Wnioski	54
5.3	Ograniczenia informacyjne	56
5.3.1	Niewiedza	56
5.3.2	Niepewność, asymetria informacyjna	57
6	Wyznaczanie rynku	60
6.1	Wprowadzenie	60
6.2	Metody i koncepcje jakościowe	61
6.3	Metody ilościowe – test SSNIP	65
6.3.1	Koncepcja podstawowa	65
6.3.2	Cellophane Fallacy („błąd celofanowy”)	68

6.3.3 Wykorzystanie testu SSNIP	69
7 Pozycja dominująca	73
7.1 Pozycja dominująca jednego przedsiębiorstwa / indywidualna pozycja dominująca	73
7.1.1 Udział w rynku	74
7.1.2 Siła finansowa, siła zasobów	78
7.1.3 Dostęp do rynków zakupu i sprzedaży	79
7.1.4 Powiązania	80
7.1.5 Bariery wejścia / potencjalna konkurencja	81
7.1.6 Konkurencja poprzez substytucję brzegową (marginalną)	83
7.1.7 Przeciwważna siła rynkowa	84
7.1.8 Faza rynku	86
7.2 Zbiorowa pozycja dominująca	86
8 Kontrola koncentracji	88
8.1 Połączenia horyzontalne	88
8.1.1 Efekty unilateralne	88
8.1.2 Efekty skoordynowane	91
8.1.3 Modele symulacyjne fuzji	92
8.2 Fuzje wertykalne i konglomeratowe	93
9 Unilateralne ograniczenia konkurencji	95
9.1 Ograniczanie dostępu do rynku	96
9.1.1 Ramy analizy	96
9.1.2 Drapieżnictwo cenowe	99
9.1.3 Systemy rabatowe	100

9.1.4	Sprzedaż wiązana i pakietowa	102
9.1.5	Odmowa dostaw	103
9.1.6	Rynki wtórne	105
9.2	Dyskryminacja cenowa	106
9.3	Zachowania eksploatacyjne	108
10	Porozumienia wielostronne	110
10.1	Porozumienia poziome / horyzontalne	110
10.1.1	Dopuszczalne porozumienia.....	111
10.1.2	Niedopuszczalne porozumienia („kartele hardcore“)	112
10.2	Porozumienia wertykalne	120
10.2.1	Założenia ogólne	120
10.2.2	Niektóre istotne cechy porozumień wertykalnych	121
10.2.3	Ocena ograniczeń wertykalnych z punktu widzenia ich wpływu na konkurencję.....	125
10.2.4	Poszczególne formy ograniczeń wertykalnych	127
	Literatura.....	130

A. Podstawy teoretyczne

1 Wprowadzenie

1.1 Jak korzystać z tego podręcznika

Niniejszy podręcznik – jako wspólna praca polsko-niemiecka – kierowany jest w szczególności do pracowników UOKiK oraz Bundeskartellamt (Federalnego Urzędu Kartelowego) i pomyślany jest jako źródło wiedzy specjalistycznej przy prowadzeniu postępowań z zakresu ochrony konkurencji. W możliwie przejrzysty sposób przedstawiono i omówiono tu możliwości praktycznego zastosowania podstawowych koncepcji i narzędzi ekonomicznych, jak również ich podstawy teoretyczne. Celem niniejszej pracy jest danie użytkownikowi do ręki narzędzia, które pomoże odpowiedzieć na konkretne pytania w zakresie prawa konkurencji i ekonomii. Ponadto staraliśmy się, aby czytelnik mógł samodzielnie ocenić celowość użycia metod analizy ekonomicznej oraz wiarygodność uzyskanych dzięki nim wyników. W tym celu zostały przedstawione punkt po punkcie, jednak bez wchodzenia w szczegóły prawa konkurencji, podstawy teoretyczne oraz możliwości zastosowania wspomnianych metod w formie przykładów wziętych z praktyki europejskiego prawa konkurencji. Aby całość zaprezentować w sposób możliwie zwarty i przejrzysty, zrezygnowano w znacznym stopniu z algebraicznego przedstawiania podstaw teoretycznych. Wybrane wskazówki dotyczące literatury pozwalają czytelnikowi na samodzielne pogłębienie wiedzy.

Celowi tego podręcznika podporządkowana została klarowna, dwuczęściowa struktura: w części A przedstawiono podstawy teoretyczne, które należy poznać celem zrozumienia części B poświęconej praktycznemu zastosowaniu instrumentów analizy ekonomicznej. Część A, oprócz ważnych, podstawowych pojęć i zasad ekonomicznych oraz podstaw neoklasycznej teorii rynku i ceny, przedstawia także najważniejsze koncepcje wykorzystywane w ekonomii przemysłu. Część B, wychodząc od głównych koncepcji polityki konkurencji (takich jak wyznaczanie rynku właściwego i siła rynkowa) przedstawia możliwości stosowania metod i koncepcji ekonomicznych w zakresie kontroli koncentracji przedsiębiorstw, nadzoru nad zachowaniami przedsiębiorstw dominujących na rynku, a także stosowania prawa konkurencji w odniesieniu do poziomych i pionowych ograniczeń konkurencji.

Wzajemne odnośniki mają ułatwić korzystanie z podręcznika i zapewnić szybką orientację w nim.

1.2 Zastosowanie metod analizy ekonomicznej w praktyce prawa konkurencji

Trudno jest obecnie wyobrazić sobie niestosowanie metod analizy ekonomicznej przy rozpatrywaniu spraw z zakresu prawa konkurencji. Temat „ekonomia i prawo konkurencji” jest jak najbardziej aktualny. Określa się go często w dyskusjach mianem "more economic approach" bądź też "more economics based approach". Jednak często nie jest do końca jasne, co należy dokładnie rozumieć pod używanym przez wiele stron pojęciem "more economic approach" („bardziej ekonomicznego podejścia”). Z pewnością nieprawdą jest, jakoby zastosowanie metod ekonomicznych w prawie konkurencji było czymś zupełnie nowym. Przykładowo niemieckie prawo konkurencji od chwili powstania w latach 50-tych bazuje na jasnych, ekonomicznych podstawach teoretycznych, a praktyka administracyjna Federalnego Urzędu Kartelowego kierowała się zawsze przesłankami ekonomicznymi. Również we wcześniejszych latach rozwój teorii ekonomii znajdował odzwierciedlenie zarówno w tworzonej prawie jak i praktyce administracyjnej urzędu. Już wtedy w wyniku rozwoju teorii ekonomii stale zmieniały się i zwiększały wymagania stawiane pracownikom organów antymonopolowych dotyczące znajomości zagadnień ekonomicznych.

Podstawową cechą "more economic approach" jest odejście od ściśle formalno-prawnych wniosków (form-based approach) w stronę indywidualnej analizy poszczególnych przypadków oraz oceny ich następstw (effects-based approach). Oznacza to na przykład, że poniżej określonego progu wartości udziału w rynku w zasadzie nie występują negatywne skutki dla konkurencji wynikające przykładowo z porozumień wertykalnych, wobec których obowiązywał dotąd bezwzględny zakaz określony w tzw. „czarnej klauzuli”.

Częściowo podejście "more economic approach" związane jest z wymaganiem, aby częściej stosować metody analizy ekonomicznej, szczególnie ilościowe, w praktyce stosowania prawa konkurencji ("zwiększona ekonomizacja"). Wymóg ten wydaje się szczególnie słuszny w obliczu szeregu korzyści wynikających z zastosowania tych metod. Niniejszy podręcznik zawiera dyskusję na temat niektórych metod i ma za

zadanie prezentację możliwości i ograniczeń ich zastosowania. W ten sposób powinien on ułatwić ocenę określonych metod i modeli oraz przyczynić się do merytorycznej debaty dotyczącej ich zastosowania.

Należy jednak pamiętać, iż oprócz zalet, metody ekonomiczne mają także swoje wady. Na przykład stosowanie metod ilościowych wymaga z reguły od organu ochrony konkurencji dużo czasu i zaangażowania pracowników (np. przy pozyskiwaniu danych). Dlatego też metody ekonomiczne nie zastępują powszechnie stosowanych analiz, ale powinny sensownie je uzupełniać. Specyficzne koszty użycia metod ekonomicznych powinny być porównane z korzyściami wynikającymi z ich zastosowania. Decyzja za lub przeciw zastosowaniu określonych metod ekonomicznych zależy od konkretnego przypadku i zawsze będzie miała charakter normatywny.

2 Analiza mikroekonomiczna

2.1 Podstawy

Mikroekonomia jest jedną z dziedzin ekonomii, którą należy wyraźnie oddzielić od makroekonomii. **Makroekonomia** przedstawia gospodarkę na poziomie wartości zagregowanych i związków całościowych. W przeciwieństwie do mikroekonomii makroekonomia posługuje się wielkościami zbiorczymi, na przykład sumą dochodów wszystkich gospodarstw domowych. Bada ona łączny stan gospodarki w takim znaczeniu, jak np. zmiana wielkości dochodów lub poziomu zatrudnienia, stopa inflacji lub wahania koniunktury.

Mikroekonomia zajmuje się z kolei zachowaniami poszczególnych podmiotów gospodarczych, takich jak pojedyncze gospodarstwa domowe i przedsiębiorstwa. Równocześnie rynki analizowane są jako (realne lub wirtualne) miejsca styku popytu i podaży w odniesieniu do jednego towaru. Mikroekonomia bada tym samym związki występujące w pojedynczych dziedzinach gospodarki. Jednostki postrzegane są jako źródło siły roboczej oraz jako konsumenci wyprodukowanych towarów, które zużywają mając na celu maksymalizację własnej użyteczności. Przedsiębiorstwa stosują czynniki produkcji, takie jak ziemia, kapitał, praca i wiedza mając na celu maksymalizację zysków.

W dziedzinie mikroekonomii **teoria konsumenta** zajmuje się stroną popytu na rynku towarowym. Ważnym przedmiotem badania jest tutaj użyteczność uzyskiwana przez konsumenta z nabytego w określonym czasie koszyka dóbr. Ważną rolę odgrywają przy tym relacje preferencyjne i krzywa obojętności¹. Naprzeciw tej teorii znajduje się **teoria produkcji**, która zajmuje się stroną podaży na rynku towarowym. Badania prowadzi się w oparciu o rzeczywistą funkcję produkcji, która prezentuje stosunek nakładów i wyników, tj. jakie ilości wyprodukowanych towarów, przy jakich nakładach mogą być wytworzone.

Istotna dla analizy ekonomicznej procesów konkurencji w ramach mikroekonomii jest **teoria rynku i ceny**, która bada wyniki zetknięcia się popytu i podaży w różnych warunkach konkurencji. W ten sposób na rynku towarowym, w zależności od jego formy rynkowej (por. r. 4.1), systematycznie występują różne ceny zbytu. **Teoria gier** (por. r. 4.5.2) rozszerza analizę o następujące po sobie w czasie interakcje większej liczby uczestników rynku (zachowanie strategiczne).

Analizy mikroekonomiczne opierają się na tzw. **indywidualizmie metodologicznym**, w ramach którego przy podejmowaniu decyzji lub w zachowaniu pojedynczej osoby bierze się pod uwagę pogląd jednostki a nie np. kolektywu osób. W ramach podstawowych założeń antropologicznych główną rolę odgrywa osoba działająca racjonalnie i egoistycznie, tzw. **homo economicus**. Myśli on w kategoriach **alternatyw** (działania), rozważa je na podstawie własnych korzyści i dokonuje wyborów, które uważa za najkorzystniejsze dla siebie (**egoizm**).

W ekonomii nie przyjmuje się jednak założenia, że decyzje podejmowane przez jednostki zawsze maksymalizują ich korzyści. **Racjonalność** rozumie się raczej w ten sposób, że podmioty gospodarcze przy pomocy posiadanych środków osiągają najlepsze rezultaty lub, że osiągają określony rezultat przy najmniejszych możliwych nakładach. Zaangażowanie środków dotyczy w szczególności intensywności poszukiwań i przetwarzania informacji, które stanowią podstawę oceny działań alternatywnych. Racjonalność w sensie ekonomicznym oznacza również, że ludzie są w stanie sklasyfikować bez sprzeczności swoje preferencje (**założenie**

¹ Krzywa obojętności jest geometrycznym przedstawieniem wszystkich kombinacji ilości dóbr, które przynoszą jednostce taką samą korzyść, jest jej więc obojętne którą z nich wybierze.

przechodniości) i dla każdej grupy dwóch dóbr mogą wykazać zdecydowane preferencje (**jednoznaczność preferencji**).

Jednostki ograniczone są w swoim działaniu przez formalne (np. ustawy) lub nieformalne (np. podstawowe wartości społeczne) granice względnie **Instytucje**. Przyjmuje się założenie, że ludzie tylko dlatego wchodzą w relacje społeczne, ponieważ rozumiane są one jako system **wymiany** usług wzajemnych, który w całości jest dla nich korzystny. Bez relacji społecznych byłoby oni wyłączeni z działań wymiennych. Wobec braku przymusu wymiana łączy się zawsze ze wzrostem korzyści netto poszczególnych uczestników. Wynika to stąd, że dwóch działających egoistycznie ludzi zgodzi się dobrowolnie na wymianę tylko wtedy, gdy wymiana taka wydaje im się korzystna, a przynajmniej neutralna. Nie oznacza to bynajmniej, że wymiana dóbr między uczestnikami rynku nie ma wpływu na jednostki w tej wymianie nieuczestniczące (ekonomia mówi w takim przypadku o „**efektach zewnętrznych**“; por. r. 5.1). Niemniej jednak pozytywne jak i negatywne zmiany poziomu użyteczności niezaangażowanych osób trzecich nie są brane pod uwagę podczas dokonywania kalkulacji wymiany przez osoby uczestniczące.

Ekonomia wychodzi w większości przypadków z założenia, że korzyść wywodząca się z konsumpcji pewnego dobra lub działalności wzrasta wraz z liczbą dóbr konsumowanych względnie z zakresem własnych poczynąń (założenie nienasylenia). Jednak zakłada się również, że wraz z konsumpcją każdej kolejnej jednostki dóbr względnie wraz z kolejnym rozszerzeniem zakresu działalności gospodarczej odczuwalny przyrost korzyści jest coraz mniejszy (lub – mówiąc inaczej – zmniejsza się **użyteczność krańcowa**). Określa się to mianem „zasady malejącej użyteczności krańcowej“ (w odniesieniu do konsumpcji) względnie jako „zasadę malejących przychodów krańcowych“ (w odniesieniu do produkcji). Jednostka tak długo będzie poszukiwać możliwości wymiany, jak długo strata korzyści płynących z rezygnacji z konsumpcji wzgl. produkcji jakiegoś dobra będzie dla niej mniejsza od przyrostu korzyści wynikających z konsumpcji lub produkcji innego dobra. Zjawisko to określa się jako prawo malejącej użyteczności krańcowej.

2.2 Modele

2.2.1 Założenia ogólne

Mikroekonomia posługuje się modelami, które tworzy się na zasadzie uproszczenia przypadków wziętych z rzeczywistości. Można wyróżnić – w zależności od tego, czy dany model ma analizować stan, zmiany czy też proces – **modele statyczne, komparatywne (porównawcze) i dynamiczne**. Dalej można podzielić modele na podstawie ich zakresu, w zależności od tego, czy badany jest jeden wyizolowany rynek (**analiza częściowa**) czy też cały układ rynkowy (**analiza ogólna**).

Podstawowe wnioski dotyczące ekonomii można przedstawić w formie **werbalnej** lub **graficznej**, bez potrzeby znajomości wyższej matematyki. Modele graficzne umożliwiają względnie szybki wgląd w ekonomiczne związki przyczynowo-skutkowe. Jednak ich przejrzystość zmniejsza się wraz z rosnącym stopniem skomplikowania, poza tym nie pozwalają na przedstawienie procesów dynamicznych. W modelach graficznych najczęściej umieszcza się na poziomej osi X skalę ilościową a na pionowej osi Y najczęściej skalę cenową. Następnie do wykresu wprowadza się linie, które przedstawiają stosunek wzajemny zmiennych umieszczonych na osiach X i Y. Następnie należy zidentyfikować punkty przecięcia i styczności powyższych krzywych. Od tych punktów prowadzi się następnie linie poziome i pionowe do osi Y lub X i odczytuje na nich wartości. W celu przeprowadzenia analizy porównawczo-statystycznej przesuwają się te linie obserwując zmiany punktów przecięcia i punktów styczności. W niniejszym kompendium stosowane będą głównie modele graficzne.

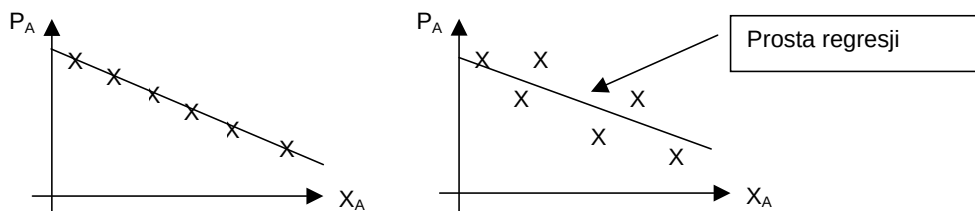
Jeżeli potrzebne są wyniki ilościowe lub doświadczalne sprawdzenie stworzonych modeli lub chcemy przeprowadzić analizę dynamiczną, nie pozostaje nic innego, jak skorzystać z metod **ekonometrycznych**. Ekspertyzy ekonomiczne, przedkładane przez przedsiębiorstwa biorące udział w postępowaniach dotyczących naruszenia prawa konkurencji lub kontroli fuzji, mają najczęściej właśnie charakter ekonometryczno-matematyczny. **Ekonometria** jest metodą, która umożliwia sprawdzenie i kwantyfikowanie teorii i związków ekonomicznych za pomocą metod statystycznych przy użyciu danych empirycznych. Jednym z najpowszechniej stosowanych instrumentów ekonometrii jest analiza regresji.

2.2.2 Analiza regresji

2.2.2.1 KONCEPCJA GŁÓWNA

Analiza regresji oznacza postępowanie, przy pomocy którego można określić związki między wieloma mierzalnymi zmiennymi. Celem analizy regresji jest wyjaśnienie określonej zmiennej przy pomocy jednej lub wielu innych zmiennych. Wychodząc od empirycznie mierzalnych danych szacuje się przy pomocy analizy regresji parametry, które umożliwiają wnioski dotyczące związków pomiędzy rozpatrywanymi danymi empirycznymi.²

Punktem wyjścia analizy regresji jest stworzenie **równania regresji**. Równanie regresji bazuje na zakładanej ekonomicznej hipotezie o związkach pomiędzy różnymi zmiennymi. Składa się ona w najprostszym przypadku ze zmiennej wyjaśnianej (**zmienna endogeniczna**) i zmiennej wyjaśniającej (**zmienna egzogeniczna**). Bada się przy tym, jaki wpływ ma zmienna wyjaśniająca na zmienną wyjaśnianą. Pomoże to zrozumieć poniższy przykład: Należy wyliczyć związek pomiędzy ilością a ceną określonego produktu, aby na tej podstawie wyliczyć np. krzywą popytu. W tym celu tworzy się pary danych ceny i ilości, tak jak to przedstawiono na ilustracji. Następnie przeprowadza się linię prostą w ten sposób, aby jak najlepiej pasowała do danych. Taką prostą nazywamy **prostą regresji**.



Wykres 1: Prosta regresji

Związek pomiędzy ceną P_A a ilością X_A jest przedstawiony na powyższej ilustracji w postaci linii prostej, przebiegającej przez punkty wyliczone na podstawie danych. Z reguły jednak nie można przyjmować, że wszystkie wyliczone punkty danych znajdować się będą na linii prostej, lecz raczej w pewnym stopniu odbiegać będą od

² Jako literaturę wprowadzającą w zagadnienia ekonometrii patrz np. Maddala i Lahiri (2001) lub Dougherty (2006).

prostej regresji. Powodem tego jest fakt, że na wyjaśnianą zmienną wpływ mają jeszcze inne czynniki, nie ujęte w modelu regresji.³ Przy pomocy procesu wyliczeń statystycznych prosta wyliczana jest w ten sposób, że punkty danych opisane w możliwie dokładnie, tzn. minimalizuje się odległość pomiędzy poszczególnymi punktami danych.

Formalnie regresja przebiega następująco. Najpierw wylicza się szacunkowe równanie regresji. W powyższym przypadku miałoby ono postać:

$$X_A = \alpha + \beta \cdot P_A + \varepsilon.$$

Parametry α i β określane są jako współczynniki. Określają one pozycję prostej regresji i są wyliczane za pomocą metod statystycznych. Wyraz ε obejmuje wpływ czynników nieuwzględnionych w równaniu, nazywa się go również składnikiem losowym.

Właściwa regresja – określenie parametrów α i β – następuje na podstawie wyliczeń komputerowych za pomocą metod statystycznych. Wynik regresji może wyglądać przykładowo jak poniższe równanie:

$$\text{Ilość A} = 2.000 - 10 \cdot \text{cena A}.$$

Mówi on nam, że przy wzroście ceny o jedną jednostkę ilość spada o 10 jednostek.

Analizę regresji można także przeprowadzić na podstawie wielu zmiennych egzogenicznych. W takim przypadku można określić łączny wpływ zmiennych oraz wpływ każdej poszczególnej zmiennej na zmienną endogeniczną. Dla przykładu można wysunąć hipotezę, że wielkość sprzedaży produktu A wyjaśnić można za pomocą cen produktów A i B. Równanie regresji w tym przypadku wyglądałoby następująco:

$$\text{Ilość A} = \alpha + \beta_1 \cdot \text{cena A} + \beta_2 \cdot \text{cena B} + \varepsilon.$$

Jako zbiór danych stosuje się przykładowo ceny miesięczne i wielkość sprzedaży produktu. Tak też w powyższym przypadku należy przyjąć założenie, że punkty danych w pewnym stopniu są rozproszone.

³ W powyższym przypadku można wyjść z założenia, że wielkość sprzedaży produktu nie jest uzależniona wyłącznie od ceny, lecz także od innych (np. sezonowych) czynników.

Przykład: Empiryczne badanie elastyczności cenowej

W ramach procesu badania fuzji przedsiębiorstw często empirycznie bada się elastyczność cenową popytu np. w celu zastosowania testu SSNIP. W takich przypadkach analiza regresji stanowi nieodzowny instrument. Za jej pomocą można wyliczyć elastyczność cenową popytu z danych empirycznych.

Punktem wyjścia jest tu stworzenie modelu regresji, który zawiera istotne czynniki, które determinują popyt na określony produkt. Chodzi tu o cenę, ale także o inne czynniki, jak np. ceny substytutów lub wydatki konsumpcyjne gospodarstw domowych itp.

Oszacowanie elastyczności cenowej popytu może napotkać na praktyczne przeszkody. Często istnieje nie jeden, lecz więcej substytutów, które należy wziąć pod uwagę w wyliczeniach. Jeżeli jednak oceniany jest model, w którym bierze się pod uwagę wiele czynników, to szacunki wpływu poszczególnych czynników są niewiarygodne i z tego powodu niewiele znaczące. Spowodowane jest to faktem, że nie można oddzielić od siebie w szacunkach różnorodnego wpływu jednego czynnika na inne zmienne.

Dalszy problem polega na tym, że zbiór danych musi być tym większy, im więcej stosuje się niezależnych zmiennych. Aby rozwiązać takie problemy, zalecane jest postępowanie stopniowe. W tym celu należy najpierw podzielić badane produkty na wiele grup towarowych, np. na produkty o wysokiej i niskiej cenie. Następnie szacuje się elastyczność cenową popytu pomiędzy produktami w zakresie jednej grupy produktów.

Można ponadto wziąć pod uwagę dalsze zmienne. W ramach badania ekonometrycznego na początku często uwzględnia się bardzo wiele zmiennych egzogenicznych i bada się stopień, w jakim wyjaśniają zmienną endogeniczną. Zmienne, dla których stwierdza się brak istotnego wpływu na zmienne endogeniczne, wyłączane są z badania w drugim jego etapie. W ostateczności otrzymuje się model regresji, który ograniczony jest do istotnych zmiennych.

2.2.2.2 WARTOŚĆ WYNIKÓW

Problem przy stosowaniu analizy regresji polega na tym, że dostarcza ona zawsze wynik liczbowy, także wtedy, gdy założenia przyjęte do wyliczeń opierają się na bezsensownych hipotezach. Również oszacowanie modelu opartego na założeniach ekonomicznych może doprowadzić do ekonomicznie bezsensownych wyników. Jeżeli przykładowo przedstawione powyżej oszacowanie funkcji popytu będzie miało w wyniku wartość dodatnią dla współczynnika β , to należałoby wyciągnąć wniosek, że popyt będzie tym większy, im bardziej wzrośnie cena. To zaprzecza jednak w przypadku normalnych produktów teorii ekonomicznej. Aby rozpoznać tego typu sprzeczności, należy zbadać wiarygodność poszczególnych współczynników. Po pierwsze należałoby przeprowadzić test poprawności. Sprawdza się w nim, czy szacowane czynniki posiadają odpowiedni znak +/- zgodnie z przyjętą hipotezą. Jeżeli tak nie jest, należy zmienić model ekonometryczny. Dla przykładu na badaną zmienną mogą mieć wpływ współczynniki, które należy włączyć do danego modelu.

Jeżeli pierwszy test poprawności wypadnie pozytywnie, można przeprowadzić dalsze testy, aby ocenić wartość wyników przeprowadzanej regresji. W tym celu stosuje się różne **miary dopasowania**. Jedną z nich jest **wartość t** (wzgl. statystyka t-Studenta). Stwierdzono już wcześniej, że model nie zawiera wszystkich zmiennych, które mają wpływ na zmienną badaną, a jedynie kilka zmiennych wybranych. Wyjaśniana zmienna nie może w związku z tym zostać wyjaśniona do końca. Szacowane współczynniki uwzględniają pewną nieścisłość. Uzyskana wartość β odpowiada wartości uśrednionej, oszacowanej na podstawie próby losowej. Dla wartości tej podaje się również dodatkowo odchylenie standardowe, względnie zakres zmienności. Z wartości odchylenia standardowego oraz z wartości szacowanej można wyliczyć wartość t. Analizując wartość t można stwierdzić, z jakim prawdopodobieństwem wartość szacowana odpowiada wartości podanej.⁴ Jako istotną przyjmuje się wartość szacowaną wtedy, gdy odpowiada ona z 90% prawdopodobieństwem wartości podanej; jako bardzo istotna uznawana jest ona wtedy, gdy odpowiada z prawdopodobieństwem 95% wartości podanej. Komputerowe programy ekonometryczne potrafią od razu odrzucić wartości szacowane statystycznie nieistotne.

Dalszą miarą dopasowania jest **wartość R^2** . Daje ona pojęcie o tym, w jaki sposób wszystkie zmienne egzogeniczne mogą wyjaśnić zmienną endogeniczną. Wartość R^2 zawiera się pomiędzy 0 i 1, należy ją interpretować jako udział procentowy, w zakresie którego zmiany zmiennej endogenicznej dadzą się wyjaśnić przez zmienne egzogeniczne. Im bardziej można wyjaśnić zmienną endogeniczną przy pomocy zmiennych egzogenicznych, tym bardziej wartość R^2 zbliża się do 1; im mniejsza jest moc wyjaśniająca danego modelu, tym bardziej wartość R^2 zbliża się do 0. Im wyższa jest wartość R^2 , tym bardziej wiarygodna jest regresja.

Wspomniane miary dopasowania, tj. test t i wartość R^2 , nie mogą być traktowane osobno. Może się zdarzyć przykładowo, że co prawda poszczególne oszacowane wartości parametrów wykazują niską wartość t a mimo to wartość R^2 jest wysoka. Podobnie regresja może dawać statystycznie istotne wartości szacowanych

⁴ Ścisłe ujmując testuje się negatywną hipotezę. Dla przykładu bada się, z jakim prawdopodobieństwem wartość szacowana odpowiada wartości zero. Jeżeli w wyniku testu otrzymamy np. prawdopodobieństwo 4%, oznacza to, że wartość szacunkowa różna jest od zera w 96%.

parametrów, które jednak nie wyjaśniają zadowalająco zmiennej wyjaśnianej (niska wartość R^2). Aby ocenić wiarygodność modelu regresji, zarówno test t jak i test R^2 muszą wypaść pozytywnie. Ale nawet jeżeli test t oraz wysoka wartość R^2 wskazują na dużą wiarygodność wyników analizy regresji, może się okazać, że wyniki są niewiarygodne lub wręcz błędne. Jedną z przyczyn może być przykładowo fakt, że poszczególne zmienne są ze sobą skorelowane (autokorelacja), lub że rozproszenie danych jest systematyczne (heteroskedatyczność). Aby rozpoznać tego typu odchylenia, przeprowadza się inne testy. Często konieczna jest modyfikacja modelu ekonometrycznego lub zastosowanie innych metod statystycznych dla oddania wartości szacunkowych.

3 Rynek

Pod pojęciem **rynek** ekonomiści rozumieją (abstrakcyjne lub konkretne) miejsce spotkania podaży i popytu. W realiach gospodarczych występują oczywiście rynki w różnych formach: I tak niektóre rynki są w dużym stopniu zorganizowane i dalece zestandaryzowane (np. rynek papierów wartościowych na giełdzie papierów wartościowych lub międzynarodowy rynek surowców na giełdzie towarowej), inne rynki charakteryzują się mniejszym stopniem zorganizowania (np. regionalny rynek samochodów używanych). Jak funkcjonują rynki? Odpowiadając na takie pytanie należy traktować podaż i popyt jako podstawowe siły rynkowe.

3.1 Strona popytu

3.1.1 Popyt indywidualny i popyt rynkowy

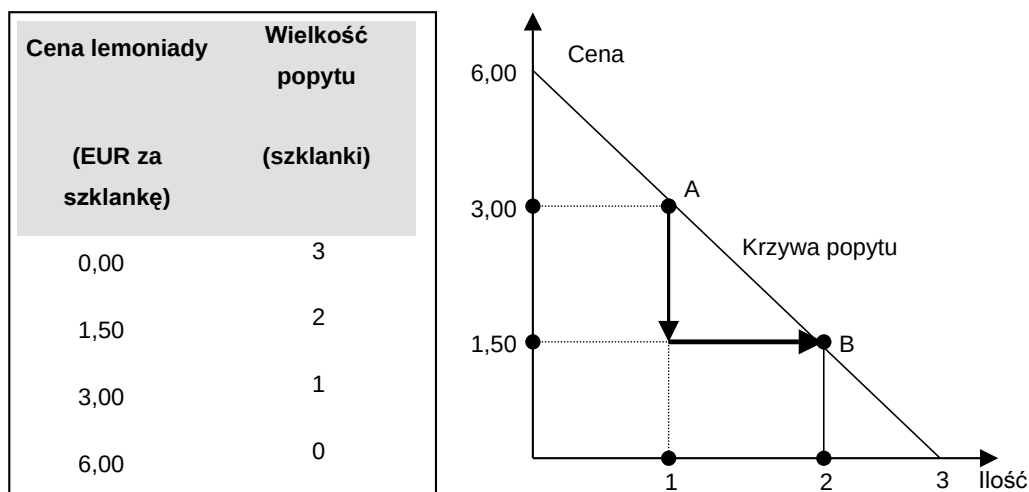
Popyt określa z jednej strony popyt jednej osoby, z drugiej strony także **popyt rynkowy** jako sumę wszystkich wielkości popytu w odniesieniu do jednego dobra lub usługi. **Popyt indywidualny** jednej osoby w odniesieniu do jednego dobra, np. lemoniady, uzależniony jest w szczególności od następujących czynników:

- **Ceny** dobra (im wyższa cena lemoniady, tym mniejszy popyt)
- **Dochodów** osoby (im wyższe dochody, tym większy popyt na lemoniadę)

- **Preferencji** (gust i upodobania) osoby (im bardziej lemoniada jest preferowana, tym większy popyt na nią)
- **Ceny innych dóbr** (im wyższa jest cena np. wody mineralnej (substytut), tym większy popyt na lemoniadę).
- Oczekiwań, czynników sezonowych itp.

Aby zbadać wpływ **zmiany ceny** na wielkość popytu, można założyć, że pozostałe wielkości, mogące mieć wpływ na wielkość indywidualnego popytu, (tymczasowo) pozostają bez zmian; mówi się w takiej sytuacji o warunku *ceteris paribus*. Przy takim założeniu można teraz stworzyć plan popytu danej osoby w odniesieniu do dobra, jakim jest lemoniada, w zależności od ceny lemoniady⁵. Plan popytu można przedstawić w formie tabelarycznej lub (graficznie) w formie diagramu.

Wykres 2: Indywidualny plan popytu i krzywa popytu w diagramie cenowo-ilościowym



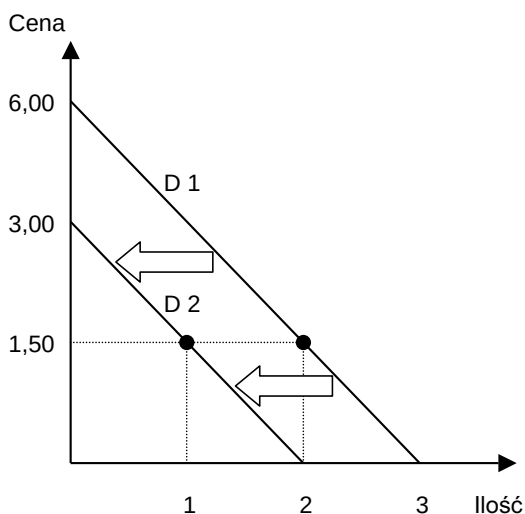
Plan popytu pokazuje nam, że przy obniżeniu ceny lemoniady z 3,00 EUR do 1,50 EUR za szklankę wielkość indywidualnego popytu wzrośnie z jednej szklanki do dwóch. Taka zmiana ceny odpowiada **przesunięciu na krzywej popytu** z punktu A do punktu B. Cena w wysokości 6,00 EUR miałaby charakter prohibicyjny, tj. osoba przestałaby konsumować lemoniadę. Opadający przebieg krzywej popytu odpowiada **prawu popytu**, według którego, przy założeniu *ceteris paribus*, dochodzi do

⁵ Cena lemoniady jest tutaj zmienną niezależną, zaś wielkość popytu na lemoniadę zmienną zależną.

zmniejszenia wielkości popytu przy rosnącej cenie danego dobra. Popyt można zaprezentować matematycznie w formie funkcji⁶; **funkcja popytu** prezentuje związek pomiędzy ceną p_i a wielkością popytu x_i na jakieś dobro.

Należy wyraźnie odróżnić przesunięcie na krzywej popytu – spowodowane zmianą ceny – od **przesunięcia krzywej popytu**. Występują one wtedy, gdy założenie niezmienności pozostałych czynników mających wpływ na wielkość popytu przestaje obowiązywać. Przesunięcie w kierunku lewym oznacza zmniejszenie, natomiast przesunięcie w prawo zwiększenie popytu przy określonej cenie. Na przykład zmniejszenie dochodów danej osoby spowoduje przesunięcie krzywej popytu w lewo (patrz. wykres 3). Zmiana preferencji, tj. kiedy np. zamiast lemoniady preferuje się piwo, implikuje również przesunięcie krzywej popytu w lewo.

Wykres 3: Przesunięcie krzywej popytu (na przykładzie spadku popytu)



Kiedy zamiast indywidualnych planów popytu badane są rynki, rozpatruje się indywidualny popyt wszystkich uczestników rynku wspólnie. **Popyt rynkowy** jest wynikiem sumy wielkości popytu wszystkich uczestników rynku deklarujących popyt na dane dobro. Graficznie tworzy się funkcję popytu rynkowego poprzez horyzontalne dodawanie do siebie wszystkich osób deklarujących popyt. Na osi

⁶ Poniżej, dla uproszczenia, prezentowane będą liniowe funkcje popytu według wzoru $x(p) = a - b \cdot p$.

ilościowej (przedstawionej horyzontalnie) otrzymujemy rosnące ilości w zależności od liczby osób deklarujących popyt.

3.1.2 Elastyczność popytu

Krzywa popytu opisuje ogólnie reakcje nabywców na rynkach, szczególnie ich reakcję na zmianę ceny. Dla stosującego prawo konkurencji pracownika organu antymonopolowego często istotna jest wiedza o tym, jak dalece wrażliwa lub elastyczna jest lub mogłaby być reakcja popytu na zmianę cen rynkowych. W pojedynczych przypadkach, w których dostępne są dane empiryczne dotyczące cen oraz związanych z nimi ilości, analiza elastyczności cenowej popytu (dokładniej: elastyczności popytu na określony produkt w odniesieniu do jego ceny) pozwala na wyciągnięcie istotnych wniosków.

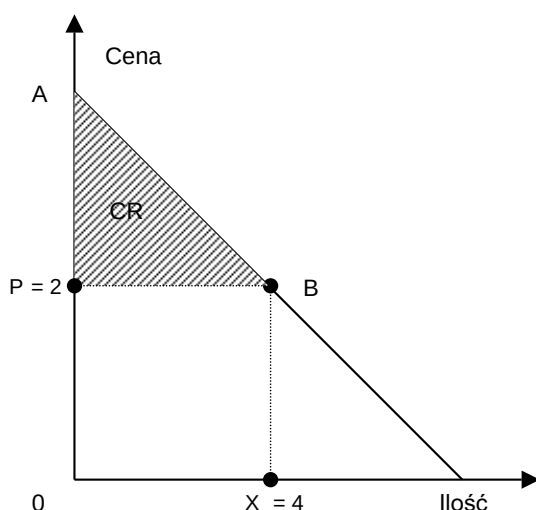
Elastyczność cenowa popytu określa, o ile procent zmieni się wielkość popytu, jeżeli cena tego produktu wzrośnie o 1 procent. Określona wartość elastyczności cenowej popytu obowiązuje zawsze wyłącznie dla jednego punktu, względnie zakresu krzywej popytu. Koncepcja elastyczności cenowej popytu jest ściśle związana z nachyleniem krzywej popytu. Obowiązuje tu ogólna zasada: im bardziej płaska (stroma) jest krzywa popytu w danym punkcie ew. zakresie, tym większa (mniejsza) jest wartość absolutna elastyczności cenowej popytu. Ponieważ nachylenie krzywej popytu w normalnym przypadku jest ujemne, także elastyczność cenowa popytu jest z reguły ujemna. Często jest ona jednak przedstawiana w wartościach absolutnych (bez znaku).

Można wyróżnić dwa (teoretyczne) przypadki ekstremalne: Popyt (w związku ze zmianą ceny) całkowicie nieelastyczny (tj. pionowa krzywa popytu), w którym to przypadku elastyczność cenowa jest równa zero, jak również całkowicie elastyczny popyt (tj. pozioma krzywa popytu), przy którym wartość elastyczności cenowej dąży do nieskończoności. Przy elastyczności cenowej popytu o wartości poniżej 1 (tzw. elastyczność jednostkowa) popyt jest nieelastyczny, przy wartościach powyżej 1 jest elastyczny. Nieelastycznym popytem charakteryzują się towary pierwszej potrzeby (np. ropa naftowa), natomiast popyt na tzw. towary luksusowe (np. drogie perfumy) z reguły jest elastyczny. Ostatecznie elastyczność cenowa popytu uzależniona jest od indywidualnych preferencji konsumentów.

3.1.3 Nadwyżka konsumenta

Ważną koncepcją przy określaniu dobrobytu ekonomicznego indywidualnych nabywców, jak również nabywców czy konsumentów łącznie, jest nadwyżka konsumenta, która wiąże się z krzywą popytu bądź z funkcją popytu. Krzywą popytu indywidualnego można również interpretować jako krzywą indywidualnej, krańcowej gotowości do płacenia. **Krańcowa gotowość do płacenia** pokazuje dla każdego dowolnego punktu na krzywej popytu, jaką cenę gotów byłby zapłacić nabywca za każdą dodatkową jednostkę danego dobra. **Nadwyżka konsumenta** to kwota, jaką kupujący *zatrzymuje dla siebie*, w sytuacji, gdy przy obowiązującym poziomie cen byłby gotowy zapłacić cenę wyższą. Związek ten można zilustrować graficznie:

Wykres 4: Nadwyżka konsumenta



Wykres pokazuje, że popyt indywidualny w przypadku ceny p_0 dla danego dobra przynosi nadwyżkę konsumenta CR w wysokości trójkąta ABp_0 . Nadwyżka konsumenta wynika z różnicy pomiędzy jego gotowością do płacenia (pole $A0X_0B$) a jego wydatkiem przy cenie p_0 (pole p_00X_0B). Gdyby nabywca był gotów zapłacić za szklankę lemoniady 3 Euro, ale cena jej wynosiłaby 2 Euro, to klient „zaoszczędziłby” 1 Euro jako nadwyżkę konsumenta. Analogicznie można określić nadwyżkę konsumenta dla wszystkich nabywców na rynku; w takim przypadku krzywa popytu odzwierciedla popyt rynkowy. Krzywa popytu rynkowego, na której do każdej ilości dobra przyporządkowana jest jedna cena, określa w ten sposób gotowość do płacenia **nabywcy krańcowego**, który jako pierwszy opuściłby rynek przy

marginalnej podwyżce ceny. Przy cenie p_0 kupować mogą wszyscy ci nabywcy, którzy wykazują gotowość do płacenia w wysokości p_0 lub wyższą. Pozostali nabywcy, którzy gotowi są zapłacić mniej niż p_0 , nie będą kupować. Obowiązuje podstawowa zasada, że im większa nadwyżka konsumenta (powierzchnia trójkąta CR), tj. ceteris paribus im niższa cena dobra, tym wyższy jest dobrobyt konsumentów.

3.2 Strona podaży

3.2.1 Podaż indywidualna i podaż rynkowa

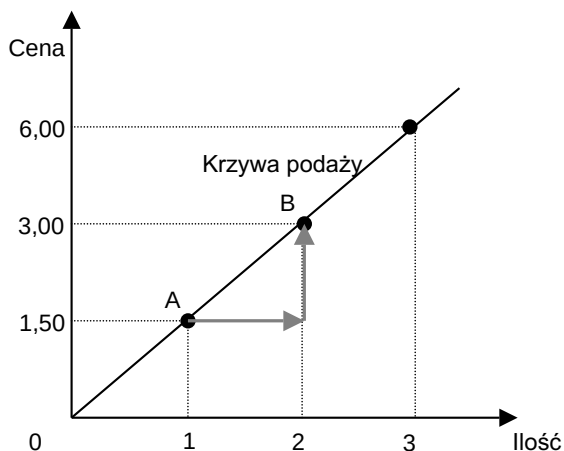
Podaż rynkowa wynika z sumy indywidualnych wielkości podaży wszystkich potencjalnych sprzedawców. Podaż rynkowa zdeterminowana jest tymi samymi czynnikami, od których uzależnione jest **zachowanie się podaży indywidualnej**:

- **Ceną sprzedaży** dobra
- Cenami **czynników** produkcji
- Dostępną technologią produkcji
- Oczekiwaniami itp.

Oferta rynkowa uzależniona jest także od liczby indywidualnych oferentów. Jeżeli wszystkie inne mające wpływ czynniki, oprócz ceny lemoniadę potraktujemy jako stałe, będziemy w stanie zbadać wpływ ceny sprzedaży na wielkość podaży. W tym celu możemy przyjrzeć się indywidualnemu planowi podaży oraz związanej z nim krzywej podaży w diagramie ilościowo-cenowym. Krzywa podaży przyporządkowuje czytelnie w układzie graficznym każdej cen dobra związaną z nią ilość podaży.

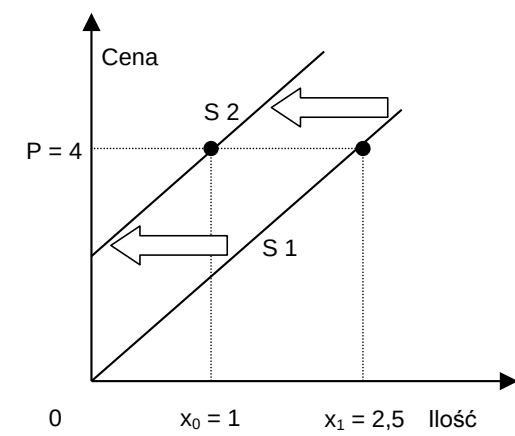
Wykres 5: Przesunięcia na krzywej podaży

Cena lemoniady (EUR za szklankę)	Wielkość podaży (Szkłanki)
0,00	0
1,50	1
3,00	2
6,00	3



Plan podaży pokazuje, że przy wzroście ceny lemoniady z 1,50 EUR do 3,00 EUR wielkość podaży wzrośnie z 1 do 2. Taka zmiana ceny odpowiada **przesunięciu na krzywej podaży** z punktu A do punktu B. Związek ten określa się także mianem **prawa podaży**, które mówi, że ilość oferowanego dobra jest tym wyższa, im wyższa jest cena tego dobra. Krzywa podaży może zostać przedstawiona także jako funkcja matematyczna⁷, która przedstawia zależność wielkości podaży i ceny pewnego dobra.

Wykres 6: Przesunięcie krzywej podaży



⁷ W poniższych przykładach przedstawiono liniowy przebieg krzywej podaży.

Należy wyraźnie odróżnić przesunięcie na krzywej podaży – spowodowane zmianą ceny któregoś dobra – od **przesunięcia krzywej podaży**. Następuje ono wtedy, gdy założenie niezmienności pozostałych czynników mających wpływ na wielkość podaży przestaje obowiązywać. Przesunięcie krzywej podaży w lewo jest konsekwencją podrożenia czynników produkcji, na przykład podwyżki cen energii elektrycznej (patrz wykres 6) Przesunięcie w prawo następuje na przykład na skutek poprawy dostępnej technologii w taki sposób, że określona wielkość produkcji może zostać zrealizowana za pomocą mniejszego nakładu kosztów.

3.2.2 Koszty i zyski przedsiębiorstwa

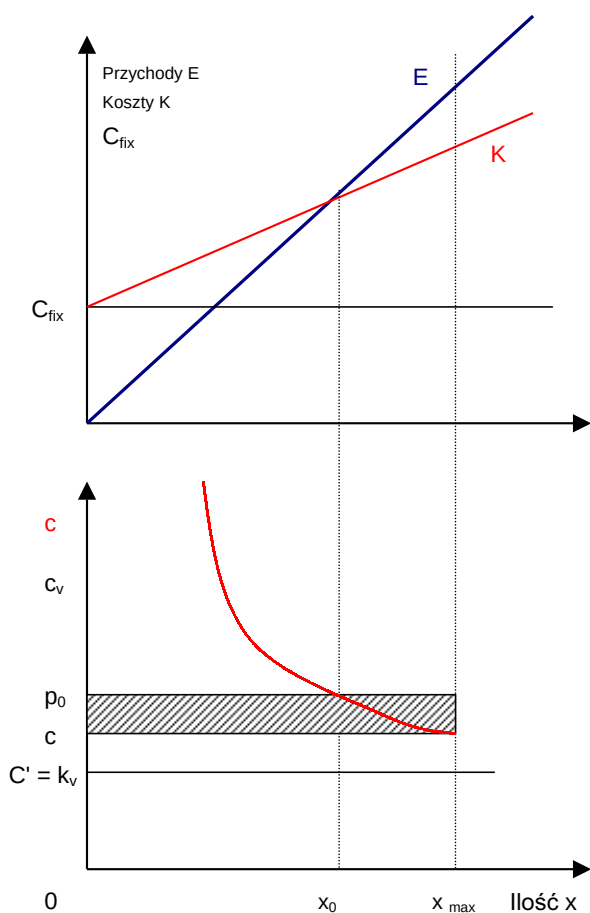
Interpretacja indywidualnych spraw rodzi konieczność znajomości istotnych pojęć z zakresu kosztów oraz ich analizy przez stosujących na co dzień przepisy prawa pracowników organów ochrony konkurencji. Dotyczy to na przykład sytuacji, w których dysponujemy poszlakami wskazującymi na próbę wyeliminowania bądź niedopuszczenia do rynku nowego operatora lotniczego przez linie lotnicze mające na tym rynku pozycję dominującą, poprzez stosowanie cen drapieżnych (tzw. predatory pricing).

Polityka kosztowa realizowana przez przedsiębiorstwa jest krokiem pośrednim na drodze do maksymalizacji zysków. Założenie teorii przedsiębiorstwa, iż podmioty z zasady dążą do maksymalizacji swoich (krótkoterminowych) zysków, nie odbiega od rzeczywistości. Jednak w podanym przykładzie drapieżnictwa cenowego założenie krótkoterminowej **maksymalizacji zysków**, o którym mowa w zdaniu poprzednim, należałoby zastąpić celem długoterminowej maksymalizacji zysków. Ogólnie obowiązuje zasada: zysk (G) równy jest przychodom (E) pomniejszonym o koszty (C). Przychody E stanowią iloczyn wyprodukowanej ilości x oraz ceny rynkowej p ($E = x \cdot p$). Koszty C są sumą iloczynów ilości i cen czynników użytych do produkcji, przy czym ceny czynników oraz cenę p należy traktować jako ustalone przez rynek. Ponieważ zarówno przychody jak i koszty przedsiębiorstwa uzależnione są od wielkości produkcji x , można przedstawić przychody i koszty jako funkcję wielkości produkcji x , względnie graficznie jako **krzywą przychodów oraz krzywą kosztów**.

Przy określonych założeniach odnośnie warunków produkcji otrzymamy w efekcie liniowy przebieg funkcji kosztów. Należy rozróżnić **koszty stałe**, czyli koszty, które nie są zależne od wielkości produkcji (np. koszty administracyjne linii lotniczych) oraz

koszty zmienne, które zmieniają się w zależności od wielkości produkcji (np. cena paliwa lotniczego). Istotne jest, aby podział na koszty stałe i koszty zmienne odnosił się do określonej jednostki czasu (np. 1 miesiąc, 1 rok). Koszty osobowe mogą być traktowane w skali jednego miesiąca jako koszty stałe w wyniku obowiązujących terminów wypowiedzenia, natomiast w skali jednego czy dwóch lat jako koszty zmienne. Koszty całkowite składają się z kosztów stałych (C_{fix}) i kosztów zmiennych (C_v). Przebieg (liniowy) krzywej kosztów i krzywej przychodów (także liniowy, ponieważ $E = p \cdot x$, gdzie p jest stałe) można przedstawić na następującym wykresie:

Wykres 7: Krzywe kosztów i przychodów



W przedstawionym przykładzie przy cenie rynkowej p_0 zysk maksymalny zostaje osiągnięty przy wielkości produkcji x_{max} , odpowiadającej maksymalnym mocom produkcyjnym przedsiębiorstwa. Zysk byłby tym większy, im więcej by

wyprodukowano. Zysk składa się z zysku jednostkowego ($p_0 - c$) przemnożonego przez maksymalną, możliwą technicznie wielkość produkcji ($x_{\max}$). Próg rentowności przypada na wielkość produkcji x_0 , przy której przychody i koszty całkowite są równe. Przy tej wielkości produkcji **średnie koszty całkowite c** , czyli koszty całkowite podzielone przez wyprodukowaną ilość produktów, odpowiadają cenie rynkowej p . Dla analizy kosztów istotnymi pojęciami są także koszty marginalne lub krańcowe C' oraz średnie koszty zmienne c_v . W przedstawionym przykładzie koszty krańcowe byłyby równe średnim kosztom zmiennym, ponieważ funkcja kosztów jest funkcją liniową o stałym nachyleniu.

Tabela 1: często stosowane pojęcia kosztowe

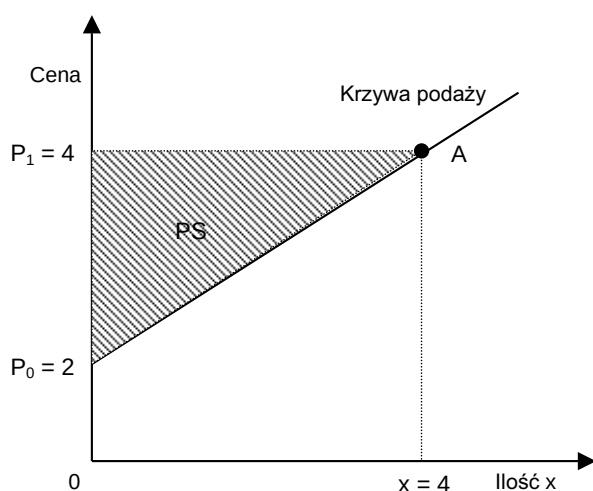
Pojęcie	Symbole	Opis
Koszty całkowite (ang. total cost)	C (C)	Suma kosztów stałych i kosztów zmiennych
(łączne) koszty stałe (ang. fixed cost)	C_{fix}	Koszty, które powstają niezależnie od wielkości produkcji i aktywności producenta
Koszty zmienne (ang. variable cost)	C_v (VC)	Koszty, które zmieniają się wraz z wielkością produkcji
Średnie koszty całkowite (ang. average total cost)	c (ATC)	Łączne koszty przypadające na jednostkę, tj. podzielone przez wyprodukowaną liczbę jednostek
Średnie koszty zmienne (ang. average variable cost)	c_v (AVC)	Koszty zmienne przypadające na jednostkę, tj. podzielone przez wyprodukowaną liczbę jednostek
Koszty krańcowe (ang. marginal cost)	C' (MC)	Koszty przypadające na wyprodukowanie kolejnej jednostki produktu (tj. pierwsza pochodna funkcji kosztów po ilości)
Średnie koszty możliwe do uniknięcia (ang. average avoidable cost)	$C_{\text{verm.}}$ (AAC)	Koszty, które nie zostałyby poniesione, gdyby określona część produkcji została zrealizowana (w przypadkach, w których można uniknąć tylko kosztów zmiennych, średnie koszty możliwe do uniknięcia odpowiadają średnim kosztom zmiennym) ⁸

⁸ Patrz: European Commission, discussion paper on the application of Article 82 of the Treaty to exclusionary abuses, op. Cit., para. 106 ff.

3.2.3 Nadwyżka producenta

Analogicznie do nadwyżki konsumenta nadwyżka producenta jest ważnym instrumentem określania dobrobytu pojedynczego oferenta oraz oferentów (producentów) łącznie. Nadwyżka producenta wiąże się ściśle z krzywą podaży. Dla indywidualnego oferenta krzywa podaży odpowiada jego krzywej kosztów krańcowych. **Nadwyżka producenta** jest różnicą pomiędzy tym, co oferent otrzyma przy istniejącym poziomie cen, kiedy odejmie od ceny sprzedaży koszty krańcowe (jego cenę oferowaną). Związek ten można dla ułatwienia przedstawić graficznie:

Wykres 9: Nadwyżka producenta



Z punktu widzenia indywidualnego oferenta wykres uzmysławia fakt, że oferent przy cenie $P_1 = 4$ oferowałby ilość 4. Różnica między kosztami krańcowymi przy danej wielkości produkcji oraz ceną określa jego nadwyżkę producenta (trójkąt P_0P_1A). Nadwyżka producenta przedstawiona jest jako powierzchnia pomiędzy krzywą podaży a linią ceny (PS).

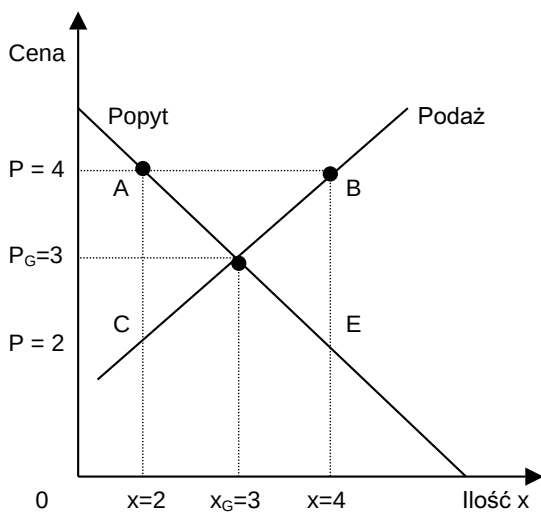
Reguła ta odnosi się także do **podaży rynkowej** oraz wszystkich producentów łącznie. Przy cenie sprzedaży $P_1 = 4$ na rynek z ofertą mogą wejść tylko ci producenci, których koszty krańcowe są niższe od ceny lub (maksymalnie) jej równe. Producenci, których koszty krańcowe są niższe niż cena sprzedaży, uzyskują nadwyżkę producenta. Wszyscy producenci, którzy wytwarzają po kosztach

krańcowych wyższych niż 4, nie będą w stanie wejść ze swoją ofertą. Producent, którego koszty krańcowe pokrywają się z ceną sprzedaży, jest **oferentem krańcowym**, który jako pierwszy opuści rynek przy obniżce ceny. Generalnie obowiązuje zasada, według której dobrobyt producenta jest tym większy, im wyższa jest cena sprzedaży i w związku z tym wyższa nadwyżka producenta.

3.3 Zależność między podażą a popytem

Równowaga rynkowa oznacza taką sytuację, w której ilość oferowana i ilość pożądana odpowiadają sobie wzajemnie. Cena, która powoduje, że ilość dóbr oferowanych i dóbr nabywanych pozostaje w równowadze, nazywana jest **ceną równowagi** lub ceną „czyszczącą” rynek, gdyż (tylko) przy takiej cenie na rynku nie powstają nadwyżki lub niedobory. Odpowiadająca temu stanowi ilość produktów to **ilość równowagi**. W stanie równowagi rynkowej siły podaży i popytu równoważą się, tj. ilość oferowana za określoną cenę odpowiada dokładnie ilości nabywanej za taką cenę. Ponieważ zarówno ilość oferowana jak i ilość nabywana pewnego dobra uzależniona jest od jego ceny, można przedstawić krzywą popytu i krzywą podaży łącznie w formie wykresu cenowo-ilościowego.

Wykres 10: Równowaga podaży i popytu



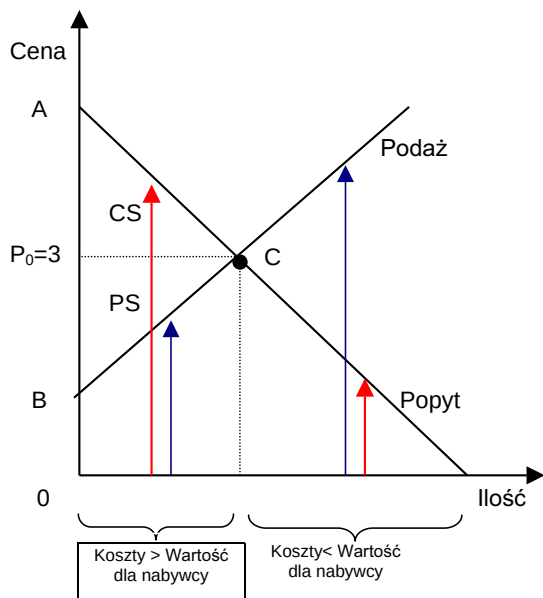
Jak działa mechanizm kształtowania się cen? Przy zbyt wysokiej cenie $p = 4$ ilość oferowanych dóbr jest co prawda względnie duża ($x = 4$), jednak nabywcy przy tym poziomie ceny nie są skłonni do nabywania dużych ilości. Stąd będą oni nabywali

względnie niewielką ilość dóbr od oferentów ($x = 2$). W następstwie tej względnie wysokiej ceny dojdzie (najpierw) do przerostu podaży (AB). Ponieważ producenci chcą sprzedać jak największą ilość dóbr, tak długo, jak długo cena sprzedaży przewyższa ich cenę oferowaną, występuje nacisk na cenę sprzedaży, która odpowiednio zaczyna spadać. To powoduje automatycznie wzrost popytu, przez co ilość oferowana i ilość nabywana zbliżają się do siebie.

Jeżeli z kolei cena sprzedaży na początku jest względnie niska, np. $p = 2$, to nabywcy chcą nabywać więcej niż oferent za taką cenę jest w stanie wyprodukować. W punktach sprzedaży tworzą się kolejki za towarem. Powstaje zatem przerost popytu (CE). Jeżeli oferenci dalej działają w ten sposób, że utrzymuje się przerost popytu, ceny zaczną stopniowo rosnąć. To powoduje automatycznie obniżenie popytu.

Posługując się oboma instrumentami – nadwyżką konsumenta i nadwyżką producenta – można ocenić różne sytuacje rynkowe, jak np. równowagę rynkową. Instrumentem do oceny stanu rynku jest **nadwyżka ogólna**, tj. suma nadwyżki konsumenta i nadwyżki producenta. Jeżeli suma nadwyżki konsumenta i nadwyżki producenta jest maksymalna, tj. kiedy stanowiąca miarę dobrobytu nadwyżka ogólna jest maksymalna, mówi się o **efektywnym wyniku rynkowym**. Można wykazać, że równowaga rynkowa, która wynika ze swobodnego działania sił podaży i popytu, jest zawsze efektywna. Graficznie można przedstawić to na diagramie cenowo-ilościowym.

Wykres 11: Nadwyżka konsumenta i nadwyżka producenta w równowadze rynkowej

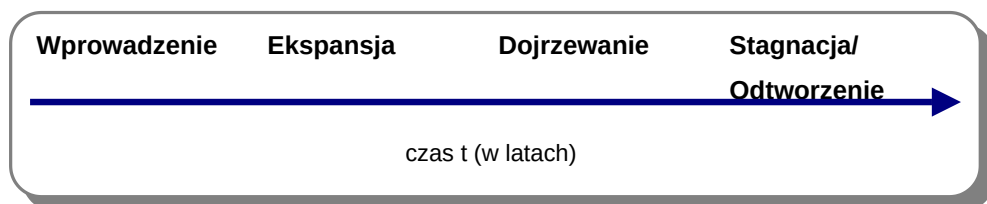


W równowadze rynkowej ($P_0; x_0$) suma wszystkich nadwyżek jest maksymalna. Nadwyżka konsumenta CR przedstawiona w formie trójkąta (AP_0C); nadwyżka producenta PR to z kolei trójkąt (BCP_0). Przy każdej założonej ilości produkcji x , która jest mniejsza od x_0 , można zwiększyć nadwyżkę ogólną w ten sposób, że zwiększy się produkcja i zbył. Ponieważ wartość dodatkowej jednostki produktu dla nabywcy znajduje się w obszarze na lewo od x_0 , będzie ona zawsze większa od kosztów wytworzenia dodatkowej jednostki dla sprzedawców. Odwrotnie przy każdej realnej ilości produkcji x większej od x_0 , można zwiększyć nadwyżkę ogólną w ten sposób, że mniej się będzie produkować i sprzedawać. Dla nabywcy wartość każdej dodatkowej wyprodukowanej jednostki produktu w obszarze na prawo od x_0 będzie zawsze mniejsza od kosztów wytworzenia dodatkowej jednostki dla sprzedawców.

3.4 Fazy rynku

Konkurencja jest procesem dynamicznym. Zgodnie z hipotezą **cyklu życia produktu** dobra stanowiące przedmiot obrotu rynkowego posiadają „cykle życiowe”⁹, w związku z czym można zaobserwować w kontekście czasowym charakterystyczne, następujące po sobie fazy rynku. Aspekt fazy rynku może odgrywać istotną rolę dla analizy konkurencyjności rynku w danym momencie wówczas, gdy faza rynku pozwala wnioskować na temat stałości sytuacji rynkowej. W określonych, pojedynczych przypadkach stadium, w jakim aktualnie znajduje się rynek, może być decydującym czynnikiem przy ocenie struktury rynku (zwłaszcza przy wysokich udziałach rynkowych) i jego perspektyw rozwoju. W procesie konkurencji można wyróżnić następujące fazy rynku:

Wykres 12: Fazy rynku



⁹ Np. przemysł motoryzacyjny w ostatnich latach charakteryzował się skróceniem cykli żywotności typów modeli.

- W **fazie wprowadzenia** powstaje nowy produkt, doprowadzany jest on do dojrzałości rynkowej i wprowadzany na rynek (tzw. innowacja produktowa, jak np. w zakresie nowych systemów telekomunikacyjnych przekazywanie obrazu (3G)). W tym stadium określanym także jako faza eksperymentalna, pozycja rynkowa i udział w rynku przedsiębiorstw określane są z reguły jako jeszcze mało stabilne i nietrwałe. Aktywność konkurencji w tej fazie jest najczęściej raczej niewielka z uwagi na monopolistyczną pozycję przedsiębiorstwa pionierskiego.
- Wprowadzenie na rynek uwieńczone sukcesem powoduje ponadprzeciętny wzrost przychodów w **fazie ekspansji**, ponieważ dołączają się nowe rzesze nabywców. Inne przedsiębiorstwa wchodzą na rynek i imitują produkt, z tego powodu na skutek rosnącej aktywności konkurencji zgodnie z oczekiwaniami cena spada. We wczesnej fazie rynku nie można wykluczyć, iż wysoki udział w rynku, względnie wiodąca pozycja przedsiębiorstwa pionierskiego na tym rynku, mogą w następnej fazie ulec osłabieniu, gdyż z reguły spodziewać się można w tym okresie nieprzeciętnego wzrostu popytu, ewentualnych innowacji produktowych oraz nowych prób wejścia na rynek. Wniosek taki wymaga jednak zbadania każdego przypadku indywidualnie. Oczekiwanie, że dominująca pozycja na rynku w fazie wprowadzania ma charakter przejściowy, wiąże się szczególnie z założeniem, że rynek jest lub będzie rynkiem otwartym na wejście konkurentów. Z tego powodu w ocenie konkurencyjności należy uwzględnić także istniejące ograniczenia dostępu (np. ochronę patentową w przypadku lekarstw) oraz wszelkie tendencje do blokowania dostępu do rynku względnie "zamknięcie rynku" (np. w przypadku oprogramowania takiego jak systemy operacyjne dla komputerów osobistych).
- Kiedy popyt zostanie w znacznym stopniu zaspokojony, następuje **faza dojrzałości**, w której zaostrza się konkurencja i powstaje nacisk na racjonalizację. Cechą charakterystyczną jest z reguły duża aktywność konkurencji. Przedsiębiorstwa mogą poprawić swoją pozycję rynkową praktycznie tylko kosztem konkurentów. W wymiarze, w jakim zmniejszają się możliwości innowacyjne, rośnie nacisk na obniżenie kosztów, a nieefektywni oferenci muszą opuścić rynek. Ponieważ warunki konkurencji w późniejszych fazach rynku zmieniają się wolniej, nie należy raczej oczekiwać spadku wysokich udziałów w

rynku. Od fazy dojrzałości następuje tendencja w kierunku rosnącej koncentracji podaży.

- W **fazie stagnacji, względnie schyłkowej** rynku następuje stagnacja lub nawet spadek podaży (np. w przypadku lodówek wyłącznie popyt odtworzeniowy). Nacisk na zmniejszenie kosztów i racjonalizację dalej wzrasta, w ten sposób prawdopodobne są (dalsze) odejścia z rynku. W przypadku jednolitych towarów masowych (np. beton towarowy) w tej fazie można zaobserwować tendencję do porozumiewania się.¹⁰ Nowe ożywienie konkurencji jest jednak możliwe, choćby wtedy, gdy dotychczasowy rynek na skutek innowacji technologicznych przemieni się w nowy rynek.

4 Formy rynku i kształtowanie cen

4.1 Założenia ogólne

Schematyczna typologia form rynkowych polega na przeciwstawieniu sobie liczby uczestników rynku po stronie oferenta z odpowiednią liczbą po stronie popytu. Tego typu podejście odgrywa na przykład ważną rolę przy ocenie równowagi sił po obu stronach rynku. Gdy spojrzymy na typowy rynek konsumentów (np. rynek detaliczny artykułów spożywczych), to mamy tam do czynienia zawsze z wieloma (małymi) nabywcami. W zależności od tego, czy naprzeciw wielu nabywców staje jeden (duży), niewielu (średnio dużych) czy też wielu (małych) oferentów, mamy do czynienia z powszechnie znanymi formami rynkowymi, takimi jak **monopol** (np. w przypadku nowo wprowadzonego leku), **oligopol** (np. w lotnictwie pasażerskim na określonych liniach, obsługiwanych przez kilku przewoźników lotniczych) względnie **polipol** (np. regionalne rynki nieruchomości). Przyjmując założenie, że wielkość uczestników rynku jest w każdej formie bardzo zbliżona, wyróżniamy dalszych sześć form rynku (patrz Wykres 13). W realiach rynkowych spotkać można m.in. (ograniczony) monopson, tak, jak ma to miejsce w przypadku aktywności popytowej

¹⁰ Podręcznikowym przykładem są zmiany na rynku transportu morskiego: kiedy około 1900 roku zbliżał się koniec okrętów żaglowych, rosnąca konkurencja ze strony okrętów parowych wywołała najpierw wojnę cenową w żegludze dalekomorskiej; zakończyła się ona dopiero po powstaniu tzw. konferencji liniowych (tj. praktycznie poprzez kartel cenowy). Kartele cenowe powstałe w ramach konferencji liniowej korzystały przez wiele lat - także w UE - ze zwolnienia z zakazu tworzenia karteli.

ze strony jedyne przedsiębiorstwa państwowego, jakim są koleje państwowe na podkłady kolejowe. Monopol bilateralny występuje wtedy, gdy jedyne przedsiębiorstwo państwowe – koleje państwowe - w danym kraju zakupują wagony od jedyne wytwórcy wagonów.

Wykres 13: Formy rynku według liczby uczestników (założenie symetryczne)

Liczba	Oferentów		
	jeden	kilku	Wielu
Nabywców			
Jeden	Monopol bilateralny	Ograniczony monopson	Monopson
Kilku	Ograniczony monopol	Bilateralny oligopol	Oligopson
Wielu	Monopol	Oligopol	Polipol

Poniżej przedstawiono model kształtowania się cen w najważniejszych formach rynku, wyróżnianych w teorii ekonomicznej rynku i cen. Jak przedstawiono powyżej, dokonuje się najpierw podziału ilościowego w zależności od liczby uczestników. Po drugie dzieli się rynki według kryteriów jakościowych, przede wszystkim na podstawie właściwości towarów stanowiących przedmiot obrotu, na rynki doskonałe i niedoskonałe. **Rynki doskonałe** charakteryzują się tym, że szczególnie spełnione są na nich następujące warunki (tzw. warunki homogeniczności):

- Z punktu widzenia nabywcy towary stanowiące przedmiot obrotu na rynku są jednorodne (dobra homogeniczne).
- Nie występują osobiste preferencje pomiędzy uczestnikami rynku.
- Nie występują różnice przestrzenne (np. brak kosztów transportu), które mogłyby uzasadniać odpowiednie preferencje.
- Nie występują różnice czasowe (np. odnośnie terminów dostaw), które mogłyby uzasadniać odpowiednie preferencje.

Na rynkach doskonałych występuje ponadto ze strony jego uczestników nieograniczona możliwość obserwacji istotnych danych rynkowych, tj. panuje pełna

przejrzystość rynku a prędkość dostosowywania się do zmieniających się warunków jest bardzo wysoka.

4.2 Konkurencja doskonała (polipol)

Model konkurencji doskonałej odpowiada prototypowi rynku doskonałego, zdefiniowanego powyżej. Teoria ekonomiczna mówi o rynku, na którym panuje konkurencja doskonała, względnie o **polipolu na rynku doskonałym**, jeżeli spełnia on następujące warunki: 1. Na rynku spotyka się bardzo wielu małych oferentów z bardzo wieloma małymi nabywcami ("polipol"); 2. Spełnione są ww. warunki homogeniczności, a w szczególności towary będące przedmiotem obrotu są homogeniczne („rynek doskonały”). W takich warunkach wszyscy oferenci i nabywcy traktują cenę rynkową jako daną, ponieważ każdy indywidualny oferent i nabywca ma zbyt mały potencjał, aby móc wpływać na cenę.

Jakie jest **zachowanie uczestników rynku** przy doskonałej konkurencji? Indywidualny nabywca nie ma możliwości zakupu towaru po cenie niższej niż cena rynkowa; stąd każdy nabywca dopasowuje ilość zakupionego towaru do obowiązującej ceny rynkowej. Tak zachowuje się np. konsument, który chce zakupić tabliczkę czekolady w sklepie samoobsługowym, który nie udziela żadnych rabatów. Każdy indywidualny oferent ma możliwość sprzedania na rynku całej ilości swojego dobra za podaną cenę; z drugiej strony za wyższą – nawet niewiele – cenę nie może on sprzedać nic. Krzywa popytu rynkowego, którą napotyka oferent w polipolu (polipolista), przebiega poziomo. Ponieważ przy konkurencji doskonałej wszyscy uczestnicy rynku akceptują podaną cenę rynkową, są oni w tym sensie cenobiorcami. Równocześnie dopasowują oni ilość zakupionego dobra do ceny rynkowej.

Jaką ilość oferuje **polipolista dążący do zmaksymalizowania zysku** na rynku doskonałym (ilość produkcji maksymalizującą zysk)? Kiedy oferent zamierza maksymalizować swoje zyski, musi przestrzegać prostej zasady. Niezależnie od formy rynku, dla osób chcących maksymalizować swoje zyski obowiązuje zasada: oferuj wielkość produkcji, przy której przychód krańcowy równy jest krańcowym kosztom (ogólny warunek maksymalizacji zysku). Ponieważ przy konkurencji doskonałej przychód krańcowy jest zawsze równy cenie rynkowej, dla wielkości produkcji maksymalizującej zysk obowiązuje specjalna zasada: koszty krańcowe równe cenie.

Jaki wynik rynkowy osiągany jest przy konkurencji doskonałej, jeżeli możliwe jest swobodne wchodzenie i odchodzenie z rynku? Nowi oferenci będą (tylko) wtedy wchodzić na rynek, gdy i jak długo dotychczasowi oferenci generują dodatnie zyski, ponieważ zyski stanowią zachętę do wejścia na rynek. Dotychczasowi oferenci generujący straty będą opuszczać rynek. Przy swobodzie wejścia na rynek i wyjścia z rynku oferenci nie będą realizować zysków długookresowo, jednak będą w stanie pokrywać swoje koszty. Zysk w wysokości zero oferent osiąga wtedy, gdy cena równa jest jego średnim kosztom całkowitym. Ponieważ w warunkach doskonałej konkurencji oferent maksymalizujący zyski przestrzega ponadto zasady „cena równa jest kosztom krańcowym“, produkuje on dalej po kosztach krańcowych, które równe są jego przeciętnym kosztom całkowitym. Ponieważ krzywa kosztów krańcowych przecina zawsze krzywą średnich kosztów całkowitych w punkcie minimum, przy długookresowej równowadze wszyscy oferenci produkują przy minimum przeciętnych kosztów całkowitych, tj. przy efektywnej wielkości przedsiębiorstwa. Obserwując cenę rynkową w długim okresie w warunkach konkurencji doskonałej należy stwierdzić, że cena rynkowa odpowiada minimum przeciętnych kosztów całkowitych. Oferenci nie osiągają zysku; przy każdej (jeszcze) niższej cenie rynkowej oferenci ponoszą straty.

4.3 Monopol

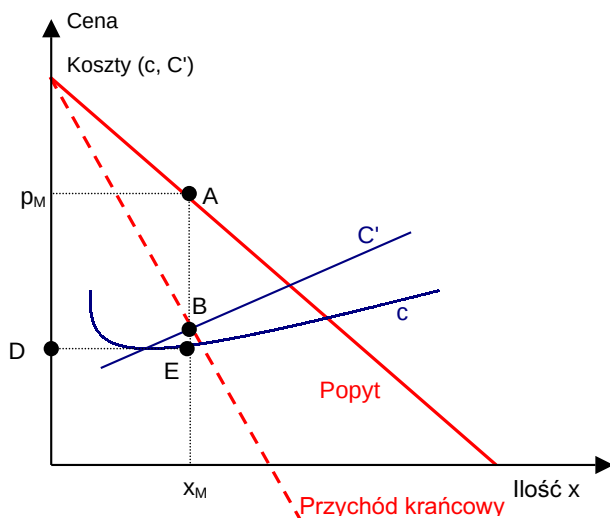
Na rynku monopolistycznym produkt oferowany jest przez tylko jednego (dużego) oferenta zwanego monopolistą. Naprzeciw niego po stronie popytowej znajduje się wielu małych nabywców. Ponieważ monopolista nie ma żadnych konkurentów, może przyjąć założenie, że nabywcy akceptują oferowaną przez niego cenę jako obowiązującą i dopasowują do tego ilość nabywanego towaru. W ten sposób monopolista dysponuje **przewagą rynkową**, wynikającą głównie z faktu, iż może on kształtować cenę na swoją korzyść. Monopole powstają w realiach rynkowych na przykład wówczas, gdy państwo udziela licencji jednemu tylko podmiotowi, który może oferować określony produkt na prawach wyłączności (np. monopol poczty).

Przewaga rynkowa monopolisty i jego potencjalny zysk ograniczone są popytem na rynku. Monopolista nie jest w stanie równocześnie narzucać ceny i ilości określonego dobra, ale decyduje tylko o jednej z tych wielkości, ponieważ wynik rynkowy przy

monopolu musi mieścić się na krzywej popytu i nawet monopolista nie zrealizuje kombinacji ilościowo-cenowej powyżej krzywej popytu.

Jaką cenę ustali monopolista? Można założyć, że monopolista jako jedyny oferent – podobnie jak oferenci w polipolu – będzie dążył do maksymalizacji zysków. Z uwagi na posiadaną przewagę rynkową nie zadowolili się jednak pokryciem średnich kosztów całkowitych i zyskiem na poziomie zera. Ponieważ zysk zawsze wynika z przychodów całkowitych pomniejszonych o koszty całkowite, monopolista będzie szukał takiej ceny, przy której różnica wynikająca z przychodów i kosztów, tzn. zysk, będzie maksymalna. Kształtowanie ceny w warunkach monopolu można przedstawić graficznie na diagramie cenowo-ilościowym.

Wykres 14: Kształtowanie ceny w monopolu



Monopolista ustala swoją cenę według (powszechnie obowiązującej) zasady kosztów krańcowych – przychodów krańcowych (warunki maksymalizacji zysku). Oznacza to, że maksymalizuje zysk przy wielkości produkcji, przy której krańcowe przychody odpowiadają kosztom krańcowym. Graficznie prezentuje to punkt przecięcia krzywej przychodów krańcowych i krzywej kosztów krańcowych. Na wykresie odpowiada to punktowi B, któremu przyporządkowana jest monopolistyczna wielkość produkcji q_M . Punkt B nie daje jednak bezpośrednio informacji o cenie monopolistycznej, przy której zysk jest maksymalny, lecz wyłącznie o monopolistycznej ilości. Wiadomo, że wynik rynkowy musi leżeć na krzywej popytu. Wielkości produkcji monopolisty q_M

może być przyporządkowany punkt A na krzywej popytu rynkowego, który na osi cenowej podaje cenę monopolistyczną.

Czym charakteryzuje się **wynik rynkowy w warunkach monopolu**? Inaczej niż w polipolu, krańcowe przychody monopolisty są zawsze niższe od ceny, ponieważ kiedy monopolista zwiększy wielkość produkcji o jedną jednostkę, zmniejsza się odpowiednio jego cena za całą sprzedaną ilość. Przychód krańcowy monopolisty może być nawet ujemny - gdyby monopolista wybrał kombinację cenowo-ilościową znajdującą się na nieelastycznym odcinku krzywej popytu. Jednak monopolista dążący do maksymalizacji zysku ustali taką cenę, aby wynik znajdował się w obrębie elastycznego odcinka krzywej popytu. Ponieważ przychód krańcowy przy maksymalnym zysku jest równy kosztom krańcowym, cena monopolisty przekracza także koszty krańcowe. Ceteris paribus cena monopolisty jest tym wyższa, im wyższe są koszty krańcowe.

Monopolista osiąga zysk. Zysk monopolisty wylicza się na podstawie zysku jednostkowego ($P - ATC$) pomnożonego przez wielkość produkcji q_M . Na wykresie zysk odzwierciedla prostokąt (AP_MDE). Przyporządkowana do niego wielkość produkcji monopolisty jest zawsze niższa niż w warunkach konkurencji doskonałej. W porównaniu z konkurencją doskonałą konsumenci otrzymują w warunkach monopolu mniejszą ilość dóbr za wyższą cenę.

4.4 Konkurencja monopolistyczna

Do tej pory omawiane były wyłącznie rynki, na których z założenia przedmiot obrotu stanowiły towary homogeniczne, czyli jednorodne. Dotyczy to w praktyce na przykład tzw. towarów masowych, jak choćby cement. Jednak istnieje wiele rynków, na których nie obowiązuje zasada jednorodności, ponadto występują **towary heterogeniczne**, w stosunku do których konsumenci wykazują zróżnicowane **preferencje**. Takie preferencje konsumentów mogą odnosić się do rodzaju i jakości wyrobów, do sytuacji w zakresie podaży w danym regionie lub osobistego preferowania określonych oferentów itp. Jeżeli rynek, tak jak ma to miejsce w przypadku konkurencji doskonałej, charakteryzuje się obecnością wielu małych oferentów i wielu małych nabywców, mówi się o konkurencji monopolistycznej lub jako synonim stosuje się określenie **polipol na rynku niedoskonałym**.

Na rynkach, na których występuje konkurencja monopolistycznej, obowiązują następujące **założenia**:

- Wielu małych oferentów staje naprzeciw wielu małych nabywców,
- Produkty heterogeniczne, w stosunku do których konsumenci wykazują zróżnicowane preferencje odnośnie jakości, aspektów czasowych, geograficznych bądź indywidualnych, a które równocześnie jako zbliżone substytuty należą do jednego i tego samego rynku właściwego produktowo,
- Poszczególni oferenci stają każdorazowo naprzeciw opadającej krzywej popytu.

W teorii ekonomicznej nie istnieje żaden powszechnie uznawany model konkurencji monopolistycznej. Istotny model do analizy takiej formy rynku stworzył CHAMBERLIN¹¹. Model CHAMBERLINA opiera się m.in. na takich założeniach, że oferenci maksymalizują swoje zyski dysponując jedynie ceną jako parametrem działań, mają wolny dostęp do rynku i są podobni do siebie nawzajem (przyjęcie założenia symetrii w odniesieniu do funkcji kosztów itd.). Istotnym efektem kształtowania cen w konkurencji monopolistycznej jest stwierdzenie: „cena produktu w konkurencji monopolistycznej jest niższa od ceny produktu heterogenicznego oferowanego w warunkach monopolu (przy założeniu, że monopolista nie posiada przewagi kosztowej), jednak wyższa, niż w warunkach konkurencji doskonałej.

4.5 Rynki oligopolistyczne

4.5.1 Ogólnie

Przedstawione dotychczas formy rynku – monopol i doskonała konkurencja – rzadko występują w realiach rynku. Najczęściej rynki charakteryzują się małą liczbą oferentów i dużą koncentracją. W takiej sytuacji jednak oferenci nie mogą podejmować swoich decyzji niezależnie, lecz muszą liczyć się z tym, że ich decyzje będą odczuwalne dla konkurentów. W tej sytuacji zachodzi pomiędzy nimi **współzależność**, ponieważ oferenci mogą reagować na decyzje konkurentów. Powstaje zatem pytanie na czym polega różnica w stosunku do efektów rynkowych powstających w warunkach monopolu lub konkurencji doskonałej. W literaturze

¹¹ Chamberlin, E. H. (1933): The Theory of Monopolistic Competition.

ekonomicznej często bada się oligopole w dwóch formach konkurencji: cenowej i ilościowej. Znane są one pod nazwą **konkurencji w sensie Cournota** (konkurencja ilościowa) i **konkurencji w sensie Bertranda** (konkurencja cenowa).

Analiza rynków oligopolistycznych jest o wiele bardziej złożona niż w przypadku monopolu lub konkurencji doskonałej z uwagi na współzależności decyzji. Każdy oferent musi podejmując decyzję liczyć się z tym, że jego konkurenci także podejmują decyzje, które uwzględniać będą podjęte przez niego decyzje. Kwestię, w jaki sposób zachowują się przedsiębiorstwa dążące do maksymalizacji zysku w tak skomplikowanej sytuacji decyzyjnej, można zbadać przy pomocy teorii gier jako instrumentu analizy decyzji strategicznych.

4.5.2 Teoria gier

Teoria gier (ang. Game Theory) przedstawia zbiór formalnych, matematycznych instrumentów analizy, przy pomocy których można przeanalizować **sytuacje decyzji strategicznych**. Modeluje się interakcje uczestników rynku w formie „gry”. Punktem wyjścia teorii gier była analiza gier prowadzonych przez przedsiębiorstwa. Coraz częściej instrumenty te zaczęto stosować w modelach z zakresu ekonomii przemysłu, ponadto także w innych dziedzinach nauk społecznych i ekonomicznych.¹²

Modele teorii gier najczęściej charakteryzują się następującymi składnikami i właściwościami:

- Pod uwagę bierze się wielu uczestników (graczy), w modelach z zakresu ekonomii przemysłu są to najczęściej oferenci.
- Każdy gracz posiada określoną ilość różnych opcji działania (strategii).
- Korzyści wzgl. zysk (wypłata) jednego z graczy uzależnione są od własnej strategii oraz strategii pozostałych graczy.
- Każdy gracz dąży do maksymalizacji swoich korzyści.

¹² Celem wprowadzenia w teorię gier patrz Binmore (1991), Myerson (1997) lub Osborne (2003). Nieco bardziej formalne przedstawienie kwestii proponuje Gibbons (1992).

Wynikająca z teorii gier analiza procesów decyzyjnych dzieli się na dwa etapy: Najpierw opisuje się strukturę decyzji oraz opcje działania poszczególnych graczy. Następnie przy pomocy instrumentów teorii gier wylicza się równowagę gry.

Analizę sytuacji decyzyjnej za pomocą teorii gier wyjaśnia poniższy przykład: Spójrzmy na wiele przedsiębiorstw konkurujących ze sobą za pomocą cen. Przedsiębiorstwa decydują o tym, czy rozpoczną wojnę cenową z konkurentem. Mogą one teraz wybrać jedną z dwóch strategii: „wysoka cena” lub „niska cena”. Jeżeli jeden gracz wprowadzi wysoką cenę a drugi niską cenę, to pierwszy poniesie stratę, a ten drugi przejmie cały popyt i uzyska wysoki zysk. Jeżeli obaj wybiorą niską cenę, to obaj osiągną jedynie niewielki zysk. Jeżeli obaj wybiorą wysoką cenę, to obaj osiągną średnie zyski. Zyski (wyплаты) przedstawione są na poniższej ilustracji. Prezentuje ona pierwszy etap – opis struktury gry.

Wykres 15: Struktura gry w dylemacie więźnia

		Przedsiębiorstwo 2	
		„wysoka cena”	„niska cena”
Przedsiębiorstwo 1	„wysoka cena”	5; 5	-5; 10
	„niska cena”	10; -5	1; 1

Zysk przedsiębiorstwa 1 przy wysokiej cenie

Zysk przedsiębiorstwa 1 przy niskiej cenie

Teraz spójrzmy na drugi etap – ustalenie równowagi. Gdyby oba przedsiębiorstwa porozumiały się co do strategii, oba wybrałyby wysoką cenę (= wynik kooperatywny). Jeżeli jednak każde przedsiębiorstwo osobno decyduje o wyborze strategii, musi się ono liczyć z decyzją konkurenta. Jeżeli na przykład przedsiębiorstwo 1 wychodzi z założenia, że przedsiębiorstwo 2 wybierze wysoką cenę, to wybierze cenę niską, ponieważ przyniesie mu to duży zysk (zysk = 10 > zysk = 5). Jeżeli przedsiębiorstwo 2 wybierze niską cenę, przedsiębiorstwo 1 także wybierze niską cenę, ponieważ daje to nadzieję na wyższy zysk (zysk = 1 > zysk = -5). Przy każdej wybranej strategii

konkurenta wybór strategii „niskiej ceny” jest dla przedsiębiorstwa 1 zawsze wyborem maksymalizującym zysk. Taką strategię nazywa się strategią dominującą. Dla oferenta 2 można przewidywać podobne posunięcia. Okazuje się, że dla niego także strategia „niskiej ceny” jest strategią dominującą. Rozstrzygnięcie gry polega na tym, że każdy gracz wybiera swoją dominującą strategię, w powyższym przypadku strategię „niskiej ceny”. Takie rozwiązanie zwane jest równowagą Nasha.¹³

Dylemat więźnia

Ta specyficzna konstelacja gry znana jest również pod nazwą „dylematu więźnia” (Prisoner's Dilemma). Bazuje ona na następującym założeniu: dwóch osadzonych w zakładzie karnym przesłuchiwanym jest osobno w kwestii popełnionego przestępstwa. Jednak nie ma żadnych dowodów w tej sprawie. Prokurator przedstawia obu następującą propozycję: jeżeli jeden z was przyzna się do winy, zostanie zwolniony jako świadek koronny, a ten drugi otrzyma najwyższy wymiar kary (= strata). Jeżeli obaj nie przyznacie się do popełnienia przestępstwa, to obaj otrzymacie niewielką karę pozbawienia wolności (= wysoki zysk). Jeżeli obaj się przyznacie, to otrzymacie obaj karę w średnim wymiarze (= niższy zysk). Więźniowie mają więc wybór pomiędzy przyznaniem albo nie przyznaniem się do winy. Jak w powyższym przypadku dominującą strategią dla każdego więźnia jest przyznanie się do winy (= niższa cena), chociaż każdy byłby w lepszej sytuacji, gdy wspólnie zdecydowali się na nie przyznawanie się do winy.

W niektórych sytuacjach decyzyjnych gra przyjmuje obrót, którego nie da się przedstawić adekwatnie w wybranej powyżej formie. Dotyczy to np. sytuacji, gdy gracze podejmują decyzje po kolei lub gdy wystąpią wielostopniowe sytuacje decyzyjne. Przy dynamicznych przebiegach gry stosuje się sekwencyjne jej przedstawienia przy pomocy drzewa struktury decyzyjnej. Decyzje pojedynczego oferenta prezentowane są w formie odgałęzień odchodzących od węzła. Rozwiązanie gry określa się w ten sposób, że dla każdego węzła decyzyjnego należy wybrać optymalną strategię dla gracza. Rozpoczyna się przy tym od czubka drzewa decyzyjnego. Takie postępowanie nazywa się również **indukcją wsteczną**

¹³ Taka koncepcja równowagi nazwana została nazwiskiem matematyka Johna Nasha.

(backward induction). Równowaga osiągana jest wtedy, gdy każdy gracz wybiera na każdym węźle decyzyjnym optymalną strategię.

Oprócz zaprezentowanych form przedstawiania i koncepcji rozwiązania gier istnieje jeszcze wiele innych. Należy brać pod uwagę w koncepcjach równowagi decyzje przypadkowe, ograniczenia informacyjne (niepewność lub niewiedza na temat decyzji innych graczy; por. r. 5.7) lub błędy decyzyjne innych graczy.

4.5.3 Model Cournota (konkurencja ilościowa)

Model Cournota opisuje sytuację rynkową, w której dwaj oferenci prowadzą walkę konkurencyjną poprzez wybór wielkości produkcji.¹⁴ Oferenci podejmują decyzję odnośnie produkowanych ilości w tym samym momencie. W przeciwieństwie do modelu monopolistycznego oferent kieruje się w swej decyzji nie tylko popytem rynkowym, ale bierze także pod uwagę fakt, że jego konkurent również wprowadzi na rynek określoną wielkość produkcji. Jego decyzja uzależniona jest w ten sposób (także) od decyzji konkurenta.

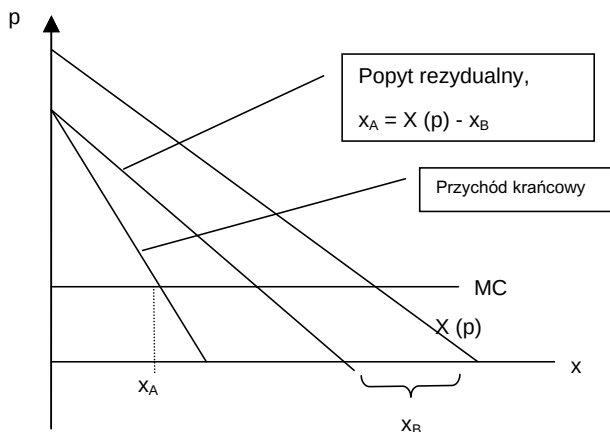
Przedstawione poniżej modele oligopolu bazują na kilku założeniach: jest dwóch oferentów A i B, którzy chcą maksymalizować swoje zyski; produkty są jednorodne; oferenci mają identyczną strukturę kosztów; nie pojawi się na rynku żaden nowy oferent i panuje na nim pełna przejrzystość. Takie założenia nie zawsze sprawdzają się w rzeczywistości, ale w znacznym stopniu upraszczają analizę.

Kalkulację decyzyjną jednego z oferentów można wyjaśnić na poniższym przykładzie. Oferent A musi brać pod uwagę, że oferent B także wprowadzi na rynek określoną ilość wyrobów. Miarodajny dla niego jest zatem popyt rezydualny, który odpowiada popytowi rynkowemu pomniejszonemu o ilość, którą oferuje oferent B: $x_A(p) = X(p) - x_B$, gdzie $X(p)$ oznacza popyt rynkowy a $x_A(p)$ i $x_B(p)$ popyt rezydualny oferenta A i B. Kalkulacja maksymalizująca zyski opiera się, podobnie jak w monopolu, na formule przychody krańcowe = koszty krańcowe. W tym przypadku dla określenia przychodów krańcowych miarodajny nie jest popyt rynkowy, lecz popyt rezydualny; odpowiada on popytowi rynkowemu pomniejszonemu o ilość, którą

¹⁴ Co prawda w rzeczywistości można często zaobserwować, że oferenci wyrobów ustalają ich cenę. Jednak do pomyślenia jest także sytuacja, w której oferent ustala swoje moce produkcyjne, a następnie je w pełni wykorzystuje.

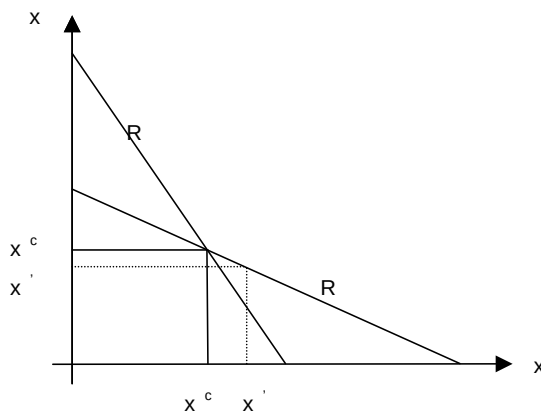
oferent B (zgodnie z oczekiwaniami) wprowadzi na rynek. Pokazuje to poniższy wykres.

Wykres 16: Popyt rynkowy w oligopolu Cournota



Jak się okazuje, popyt rezydualny dla oferenta A jest tym mniejszy, im większa jest ilość wprowadzana przez oferenta B (krzywa popytu rezydualnego przesunie się do środka), lub odwrotnie: im mniej oferent B wprowadzi na rynek, tym więcej oferent A może na nim sprzedać. Taka sama kalkulacja obowiązuje dla oferenta B: Im większa wielkość oferty oferenta A, tym mniej może on sam sprzedać. Z kalkulacji maksymalizującej zyski można wyprowadzić tak zwaną **krzywą reakcji**, która mówi, jaką wielkością produkcji „zareaguje” określony oferent na decyzję swojego konkurenta dotyczącą wielkości produkcji. Krzywa reakcji ma zawsze nachylenie ujemne względem wielkości produkcji konkurenta. Oznacza to, że wielkości produkcji są **strategicznymi substytutami (strategic substitutes)**. Funkcje reakcji przedstawione zostały na poniższej ilustracji.

Wykres 17: Krzywa reakcji w oligopolu Cournota



Kłopot oferentów przy podejmowaniu decyzji dotyczących wielkości produkcji polega na tym, że muszą oni podjąć decyzję nie znając jeszcze decyzji swojego konkurenta. Muszą zatem założyć pewne oczekiwania co do decyzji konkurenta. Przy określaniu oczekiwań pomagają im funkcje reakcji. Ponieważ panuje przejrzystość na rynku, oferenci znają **funkcję reakcji** swego konkurenta. Na tej podstawie oferent A może wyliczyć, dla dowolnej wielkości produkcji, jaką ilość zaoferuje konkurent B. Jeżeli zaoferuje ilość x_A' , musi się liczyć z tym, że oferent B zaoferuje na rynku ilość x_B' . W tym przypadku wynik rynkowy dla oferenta A nie maksymalizowałby jego zysku, ponieważ jego ilość towaru x_A' nie leży na krzywej reakcji, x_A' nie jest więc optymalną reakcją na x_B' . Takiej samej analizy można dokonać dla oferenta B. W konsekwencji oferent zachowuje się w sposób maksymalizujący zysk wybierając taką wielkość produkcji, na którą konkurent zareaguje optymalną dla niego wielkością produkcji, tj. jego wielkość/ilość produkcji znajdzie się na obu funkcjach reakcji. Równowaga rynkowa znajduje się w miejscu przecięcia obu funkcji reakcji przy wielkościach Cournota x_A^c i x_B^c .

Można matematycznie wyliczyć, że oferowana całkowita wielkość produkcji w konkurencji Cournota odpowiada dokładnie 66% wielkości podaży przy doskonałej konkurencji. Zaopatrzenie rynku w konkurencji Cournota jest zatem lepsze niż w warunkach monopolu, gdzie zaopatrzenie rynku wynosi około 50% wielkości oferowanej w warunkach konkurencji doskonałej. Cena w modelu Cournota jest w ten sposób niższa niż w cena monopolistyczna.¹⁵ Występuje jednak strata dobrobytu społecznego, ponieważ rynek zaopatrywany jest gorzej niż w warunkach doskonałej konkurencji. Jeżeli zwiększymy liczbę oferentów, to wzrośnie w ten sposób wielkość równowagi a cena spadnie. Im większa liczba oferentów, tym bardziej wynik rynku zbliża się do wyniku przy doskonałej konkurencji.

¹⁵ Fakt, że cena w modelu Cournota jest niższa od ceny monopolistycznej, daje się wyjaśnić na podstawie wykresu 15. W monopolu ilość towaru ustala się na podstawie punktu przecięcia krzywych kosztów krańcowych i przychodów krańcowych. Odpowiadająca ilości cena wynika z krzywej popytu rynkowego. W modelu Cournota wielkość produkcji poszczególnych oferentów określa się na podstawie popytu rezydualnego względnie odpowiadającej jej krzywej przychodu krańcowego. Odpowiednia cena równowagi wynika także z krzywej popytu rezydualnego. Przebiega ona zawsze poniżej krzywej popytu rynkowego, tak więc i cena, jaką można osiągnąć, znajduje się poniżej.

4.5.4 Model Bertranda (konkurencja cenowa)

Często zdarza się, że przedsiębiorstwo nie decyduje o wielkości produkcji tylko o cenie wyrobu. Jeżeli oferent ustali cenę, to jest gotów sprzedać taką liczbę wyrobów, na którą nabywcy zgłoszą popyt. Podobnie jak w modelu Cournota pojedynczy oferent nie może autonomicznie ustalać ceny, lecz musi brać pod uwagę ceny ustalane przez swoich konkurentów. Ponieważ oferenci nie znają decyzji swoich konkurentów, muszą oni zakładać pewne oczekiwania co do ich decyzji.

Kalkulację decyzyjną pojedynczego oferenta można najprościej przedstawić na przykładzie z dwoma oferentami (duopol). Przyjrzyjmy się oferentowi A. Ustala on swoją cenę p_A . Jego decyzja jest jednak uzależniona od tego, jaką cenę ustali oferent B. Pozostają mu trzy możliwości decyzji:

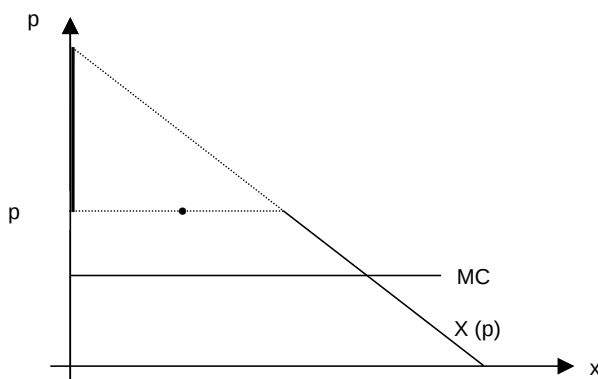
- Oferent A ustali wyższą cenę niż B: $p_A > p_B$.
- Oferent A ustali taką samą cenę: $p_A = p_B$.
- Oferent A ustali niższą cenę: $p_A < p_B$.

Jeżeli produkty są jednorodne, to klienci będą kupować u tego oferenta, który oferuje niższą cenę. Odpowiednio do tego dla oferenta A pozostają następujące wielkości zbytu:

- $p_A > p_B$: wszyscy klienci kupują u B; $x_A = 0$.
- $p_A = p_B$: oferenci A i B dzielą między siebie popyt rynkowy.
- $p_A < p_B$: wszyscy klienci kupują u A; x_A odpowiednio do popytu rynkowego.

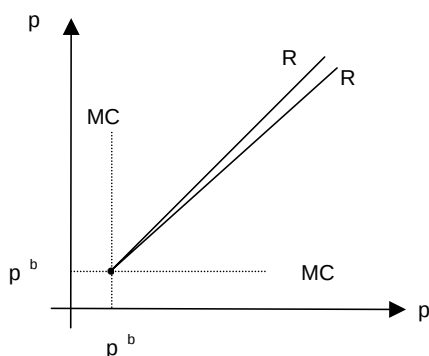
Poniższy wykres przedstawia przebieg „funkcji popytu rezydualnego” oferenta A.

Wykres 18: Funkcja popytu rezydualnego w oligopolu Bertranda



Optymalna reakcja oferenta A na cenę B da się określić na podstawie prostego rozumowania: jeżeli oferent A ustali wyższą cenę, to nic nie sprzeda i nie osiągnie zysku. Jeżeli ustali taką samą cenę jak oferent B, to podzielią między siebie popyt i osiągną taki sam zysk. Oferent A może jednak ustalić cenę, która będzie nieznacznie niższa od ceny p_B . W takim przypadku przypadnie na niego cały popyt rynkowy. Jego zysk jest prawie dwa razy większy, ponieważ obniżył on swoją cenę jedynie marginalnie, ale za to podwoił ilość sprzedawanego towaru. Obniżenie ceny jest w ten sposób zabiegiem maksymalizującym zysk. Optymalną reakcją oferenta B na cenę p_A będzie także obniżenie ceny. Reakcja oferentów jest w ten sposób skierowana w tym samym kierunku. Przyporządkowane temu funkcje reakcji oferentów przedstawiono na poniższym wykresie. Wykazują one nachylenie dodatnie. Ceny są zmiennymi **strategicznie komplementarnymi (strategic complements)**.

Wykres 19: Funkcje reakcji w oligopolu Bertranda



Jaką cenę powinien w tej sytuacji wybrać oferent A, jeżeli musi się liczyć z tym, że oferent B ustali jeszcze niższą cenę? Oferent B dopiero wtedy nie będzie mógł zaoferować niższej ceny, gdy oferent A ustali cenę na poziomie kosztów krańcowych $p_A = MC$; podanie niższej ceny przez B oznaczałoby dla niego stratę. Ponieważ dla oferenta B obowiązuje taka sama kalkulacja, wybierze on również cenę na poziomie kosztów krańcowych $p_B = MC$. W konsekwencji obaj oferenci wybiorą cenę na poziomie kosztów krańcowych $p = MC$. Całkowita zaoferowana ilość produktu odpowiada ilości konkurencyjnej, nie mamy do czynienia ze stratą dobrobytu

społecznego. Wynik rynkowy jest w ten sposób optymalny z punktu widzenia dobrobytu.

Wynik wydaje się z początku zaskakujący: chociaż jest tu tylko dwóch oferentów, to konkurencja jest tak silna, że obaj będą się przebijać ceną tak długo, aż osiągną poziom kosztów krańcowych. Zaprzecza to koncepcji, jakoby niewielka liczba oferentów sprzyjała niskiej intensywności konkurencji. Mówi się w związku z tym o **paradoksie Bertranda**.

Ten pozornie paradoksalny wynik opiera się na kilku wyżej wymienionych założeniach. Przyjęto, że oferenci produkują przy **identycznych kosztach** i posiadają nieograniczone **zdolności produkcyjne**. Jeżeli odejdziemy od tych założeń, zmieni się nasz wynik. Jeżeli oferenci mają np. różne koszty krańcowe, to ci oferenci będą jak poprzednio przebijać się nawzajem w konkurencji cenowej. Proces ustalania coraz niższej ceny kończy się jednak wtedy, gdy cena odpowiada kosztom krańcowym jednego z oferentów. Nie zareaguje on już na kolejną obniżkę ceny przez konkurenta. Oferent o niższych kosztach krańcowych ustali nieco niższą cenę niż koszty krańcowe konkurenta. W ten sposób jednak cena jest wyższa od kosztów krańcowych tańszego oferenta, a co za tym idzie, wyższa od ceny konkurencyjnej. W przeciwieństwie do konkurencji Bertranda z kosztami symetrycznymi będziemy tu mieli do czynienia ze stratą dobrobytu. Także przy ograniczeniach mocy produkcyjnych wynik Bertranda nie sprawdza się. W takim przypadku oferent od pewnej granicy nie jest w stanie obsłużyć wszystkich nabywców, jeżeli ustalił on swoją cenę poniżej ceny konkurenta. W ten sposób znajdują się nabywcy, którzy nie mogą uzyskać produktu. Jednak ci nabywcy mogą zostać zaopatrzeni przez konkurenta po cenie wyższej niż koszty krańcowe. Nie mają oni alternatywy, tak więc konkurent może ustalić cenę przewyższającą jego koszty krańcowe. Dalsze krytyczne założenie polega na tym, że chodzi w takim przypadku o grę jednookresową. Jeżeli taka sytuacja decyzyjna powtórzy się kilkakrotnie, to zmieni się sytuacja decyzyjna. W takim przypadku istnieje możliwość, że obaj oferenci wybiorą cenę wyższą niż koszty krańcowe, jeżeli mogą liczyć się z tym, że także ich konkurent w przyszłości podniesie swoją cenę.

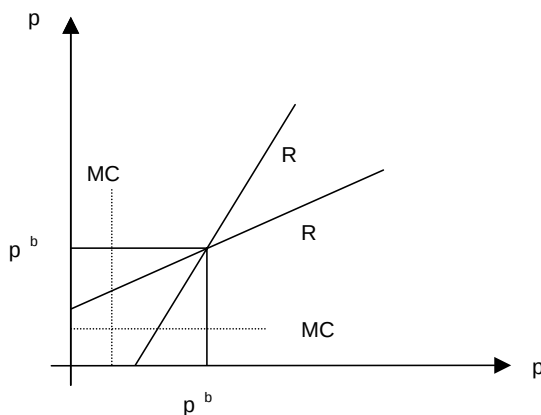
4.5.5 Rynki niejednorodne (heterogeniczne)

Na pewno najbardziej restrykcyjnym założeniem wcześniej prezentowanych modeli Cournota i Bertranda jest założenie jednorodności towaru. Takie założenie w rzeczywistości jest prawie nie do spełnienia. Najczęściej produkty różnią się od siebie, a nabywcy posiadają różne preferencje. Jeżeli przyjmiemy założenie, że na rynku występują heterogeniczne produkty, to efekty rynkowe będą różniły się od tych na rynku jednorodnym. Największa zmiana w odniesieniu do wyniku rynku zajdzie w modelu Bertranda.

Przyjęcie założenia produktów homogenicznych w modelu Bertranda ma takie następstwa, że klient zawsze kupuje u oferenta, który oferuje najkorzystniejszą cenę. W wyniku tego oferenci ustalają coraz niższe ceny aż do momentu, gdy cena odpowiada kosztom krańcowym (a dalsze obniżanie ceny nie jest już możliwe). Zmienia się to, gdy produkty są niejednorodne, a klienci preferują jeden z produktów. Jeżeli oferent zaproponuje cenę, która będzie niższa od ceny konkurenta, nie przejmie on całego popytu rynkowego, a tylko jego określoną część. Jeżeli oferent zaproponuje wyższą cenę niż jego konkurent, to nie straci on całego popytu rynkowego, a tylko jego określoną część. Spadek lub wzrost liczby klientów jest tym większy, im większa jest różnica ceny pomiędzy produktami i im mniejsza jest siła preferencji. Jeżeli klienci są względnie mało wybredni przy wyborze obu produktów, to część klientów, która zareaguje na różnicę w cenie, będzie liczna.

Z kalkulacji decyzyjnej oferenta można ponownie wyprowadzić funkcję reakcji, która przedstawiona jest na poniższej ilustracji. Cena, którą ustala oferent, uzależniona jest ponownie do ceny jego konkurenta. Im niższa jest cena p_B , tym niższą cenę ustali oferent A - p_A . Jednak obniżając cenę nie przyciągnie on do siebie całego popytu rynkowego. Nawet kiedy zaoferuje najniższą z możliwych cen, cenę równą kosztom krańcowym ($p_A = MC$), nie przejmie sam całego popytu rynkowego. Zachęta do obniżania ceny jest w tej sytuacji mniejsza. Jeżeli oferent B podwyższy swoją cenę, oferent A będzie miał również zachętę do podniesienia swojej ceny. Reakcje oferentów są więc ukierunkowane tak samo, funkcja reakcji posiada dodatnie nachylenie, nie jest jednak tak stroma, jak w przypadku funkcji reakcji przy produktach jednorodnych.

Wykres 20: Funkcja reakcji w oligopolu heterogenicznym



Punkt równowagi rynkowej znajduje się w punkcie przecięcia funkcji reakcji. Tylko w tym przypadku ceny dla obu dostawców maksymalizują ich zyski. Żaden oferent nie ma bodźca, aby ustalić inną cenę. W ten sposób **ceny** oferentów zawsze są wyższe niż **koszty krańcowe**.¹⁶ Podobnie jak w homogenicznym modelu Cournota konsekwencją tego jest spadek dobrobytu, ponieważ zaopatrzenie rynku jest gorsze niż w warunkach konkurencji doskonałej.¹⁷

4.5.6. Dalsze założenia z zakresu ekonomii przemysłu

Przedstawione powyżej modele opisują w istocie decyzje przedsiębiorstw dotyczące ilości lub ceny. W rzeczywistości jednak przedsiębiorstwo decyduje o wielu innych parametrach, które mają wpływ na działania konkurencyjne. Od lat pięćdziesiątych XX. wieku tworzone są coraz to nowe modele, przy pomocy których można zbadać dalsze parametry konkurencji, takie jak zróżnicowanie produktów, innowacje lub reklama, ale także elementy struktury takie jak bariery dostępu do rynku lub asymetria informacji. Dla takiego podejścia przyjęło się określenie **ekonomia**

¹⁶ To, że ceny w heterogenicznym modelu Bertranda są wyższe od kosztów krańcowych, można udowodnić matematycznie.

¹⁷ Także w modelu Cournota zmienia się wynik rynkowy, jeżeli przyjmimy heterogeniczność produktów, jednak nie odbiega on znacznie od uzyskiwanego na rynku homogenicznym. Cena równowagi i oferowane ilości uzależnione są od stopnia heterogeniczności produktów. Jeżeli porównamy z kolei wynik heterogenicznego modelu Cournota z wynikiem heterogenicznego modelu Bertranda, to i tu się okaże, że ceny w modelu Bertranda oraz zaopatrzenie rynku są lepsze niż w modelu Cournota.

przemysłu. O ile na początku modele te bazowały na metodach analizy mikroekonomicznej, coraz częściej zaczęto stosować teorię gier jako instrument analizy tendencji w ekonomii przemysłu. Wspólne dla modeli z zakresu ekonomii przemysłu jest to, iż zajmują się one najczęściej rynkami skoncentrowanymi z niewielką ilością oferentów. Powodem tego jest fakt, że tylko na takich rynkach zachodzi interakcja pomiędzy ich uczestnikami.

Niektóre z założeń w ekonomii przemysłu mają znaczenie dla polityki konkurencji. Bada się przykładowo, w jakich warunkach przedsiębiorstwa skłaniają się do koordynowania zachowań konkurencyjnych, jakie czynniki sprzyjają takim zachowaniom, kiedy opłaca się przedsiębiorstwu wyrugować konkurenta z rynku, jakie strategie należy zastosować, by osiągnąć taki rezultat itp.¹⁸ Zastosowanie powyższej teorii w przypadku zmów kartelowych przedstawione zostanie w rozdziale 10.1.2.2.

5 Zawodność rynku

Wolna gra sił rynkowych nie zawsze prowadzi do efektywnej alokacji zasobów. Czasami mamy do czynienia z tzw. „zawodnością rynku”. Oznacza to, że rynki załamują się, w ogóle nie powstają, są niewystarczająco wykorzystywane lub nie biorą pod uwagę wszystkich konsekwencji wiążących się z transakcjami rynkowymi pomiędzy uczestnikami rynku lub pomiędzy nimi a osobami trzecimi. Ważnym powodem są tutaj efekty zewnętrzne, niepodzielność, ograniczenia informacyjne, brak zdolności dopasowania i niepewność. Przy badaniu konkurencyjności rynku należy uwzględnić i ocenić te aspekty, aby uniknąć nieuzasadnionych interwencji, które mogłyby zaszkodzić funkcjonowaniu rynku.

¹⁸ Dalsze typowe, badane metodami ekonomii przemysłu kwestie to choćby: w jakich warunkach przedsiębiorstwa mają skłonność do różnicowania swoich produktów w strukturach poziomych lub wertykalnych? W jakich warunkach rynkowych przedsiębiorstwa mają silną zachętę do innowacji? Czy zgłaszają one niezwykłą liczbę patentów lub nie wprowadzają na rynek produktów opatentowanych, aby powstrzymać konkurentów przed wejściem na rynek? W jakich warunkach przedsiębiorstwa stosują za mało lub za dużo reklamy? Czy przedsiębiorstwa poprzez swoje zachowanie mogą stworzyć strategiczne bariery wejścia na rynek?

5.1 Efekty zewnętrzne

Efekty zewnętrzne pojawiają się zawsze wtedy, gdy występuje bezpośredni związek pomiędzy funkcjami zysku, względnie użyteczności wielu podmiotów gospodarczych, którego nie można do końca kontrolować i który nie znajduje odzwierciedlenia w relacjach rynkowych.¹⁹ Następstwem tego jest fakt, że jednostka nie ponosi sama wszystkich spowodowanych przez siebie kosztów (np. zanieczyszczenie środowiska), względnie nie otrzymuje wynagrodzenia za wszystkie oferowane korzyści. Różnicę pomiędzy kosztami całkowitymi, powstałymi dla społeczeństwa na skutek działań poszczególnych uczestników rynku a kosztami powstającymi bezpośrednio u sprawcy określa się jako efekty zewnętrzne. W przypadku negatywnych efektów zewnętrznych wielkość produkowana przez sprawcę (czyniącego szkodę) jest zbyt duża a cena zbyt niska z punktu widzenia gospodarki narodowej. I dokładnie odwrotnie dzieje się w przypadku występowania pozytywnych efektów zewnętrznych (np. badania podstawowe). Tutaj, za wygórowaną z punktu widzenia gospodarki narodowej cenę, produkuje się nieadekwatnie mało. Dochodzi do nieoptymalnej alokacji czynników, które mogą uzasadniać regulacyjną ingerencję państwa. Przyczyną efektów zewnętrznych często jest fakt, iż nie można innych uczestników rynku wyłączyć z udziału w niektórych kosztach, względnie korzyściach, które wiążą się z działaniem ekonomicznym danego podmiotu. Występowanie czynników zewnętrznych wiąże się ściśle z obowiązującą formą praw własności i kosztami ich wykonywania.

Ważnym przykładem z punktu widzenia prawa konkurencji są prace badawczo-rozwojowe. Ponieważ w warunkach braku ochrony patentowej własność intelektualna nie byłaby dostatecznie chroniona i nie występowałaby zachęta do innowacji, konkurencja na poziomie produktów zostaje na jakiś ograniczona po to, aby mogła się w ogóle rozwinąć na płaszczyźnie idei. Kiedy wynalazki stają się innowacjami, czyli pomysły przeradzają się w produkty, zwiększa się konkurencja na rynku tych produktów. Konsolidacja silnej pozycji rynkowej zostaje w ten sposób powstrzymana poprzez działania na pierwszy rzut oka ograniczające konkurencję.

¹⁹ Kiedy z powodu wzrostu popytu na płaskie monitory spada popyt (i cena) na monitory kineskopowe, mówi się wtedy o pieniężnych efektach zewnętrznych. Takie pośrednie związki nie są zawodnością rynku, lecz – wręcz przeciwnie – są konsekwencją dobrze funkcjonującego rynku.

Skrajnym przypadkiem efektów zewnętrznych są takie dobra, które z powodu braku rywalizacji w ich konsumpcji (konsumpcja dobra przez jedną osobę nie zmniejsza możliwości jego konsumpcji ze strony innych osób) oraz niemożności wyłączenia jednostek z ich konsumpcji nie mogłyby być oferowane na rynku i dlatego – o ile w ogóle – oferowane są tylko przez instytucje publiczne. Z tego powodu nazywane są też **dobrami publicznymi**. Przykładami mogą być obrona narodowa lub ochrona przeciwpowodziowa. W tych przypadkach aspiracje organu antymonopolowego w kierunku poprawy strukturalnych warunków konkurencji byłyby skazane na porażkę.

5.2 Niepodzielność

Baumol, W. (1977): On the Proper Cost Test for Natural Monopoly in a Multiproduct Industry. In: American Economic Review, Vol. 67, str. 809-822.. Train, K.E. (1995): Optimal Regulation – The Economic Theory of Natural Monopoly. Cambridge, str. 1-17.

5.2.1 Założenia ogólne

Niepodzielność dóbr lub czynników produkcji najczęściej bierze się stąd, że zdolność wykorzystania określonych zasobów może z przyczyn technicznych być zmieniana wyłącznie skokowo. Przykładem mogą tu stanowić szlaki komunikacyjne i elektrownie. Często niepodzielność prowadzi do koncentracji tej strony rynku, po której występuje. W skrajnym przypadku popyt może być zaspokajany po najniższej cenie przez tylko jednego oferenta. Mówi się wtedy o „**monopolu naturalnym**”. Występuje on w wielu dziedzinach wiążących się z zaopatrzeniem poprzez instalacje sieciowe. Monopole naturalne kryją w sobie niebezpieczeństwo, że przeciwna strona rynku znajdzie się w gorszej sytuacji z uwagi na niską jakość, zawyżone ceny lub niewystarczającą wielkość oferty, problem mogą również stanowić słabe bodźce do unowocześniania technologii.

Z reguły niepodzielność nie prowadzi sama z siebie do naturalnych monopolu, lecz do sytuacji, w której na rynku utrzyma się jedynie niewielu uczestników (oligopol). W takich przypadkach występuje niebezpieczeństwo, że oferenci będą usiłowali ograniczyć konkurencję porozumiewając się ze sobą, co doprowadzi do takich samych problemów, jak w przypadku monopolu naturalnego. Dlatego w praktyce, tak jak ma to miejsce w przypadku kontroli koncentracji, dużą wagę przywiązuje się do kwestii zbiorowej czyli kolektywnej pozycji dominującej (por. rozdział 7.2) oraz do zrozumienia przyczyn zawodności rynku skutek wyniku niepodzielności.

Przyczyn niepodzielności należy szukać w **subaddytywności**²⁰ struktur kosztów, które wiążą się z korzyściami skali lub zakresu przy produkcji jednego lub wielu dóbr.

5.2.2 Korzyści skali (economies of scale)

Korzyści skali (economies of scale) są wynikiem spadających kosztów średnich, względnie rosnących przychodów ze skali.²¹ Mogą one mieć różne przyczyny:

- Minimalne ilości wykorzystywanych czynników produkcji: Przy zwiększeniu stopnia wykorzystania mocy produkcyjnych koszty czynnika produkcji dzielą się na większą liczbę jednostek (**spadek kosztów stałych**).
- **Korzyści wielkiej partii produkcji**: jeżeli za pomocą jednego czynnika produkcji można wyprodukować więcej produktów a zmiana technologii czy sposobu produkcji wiąże się z dużymi kosztami, to opłaca się wyprodukować jak największą liczbę w jednej partii, ponieważ koszty zmian technologicznych rozkładają się na liczbę sztuk w partii i wraz z wielkością partii tracą na znaczeniu.
- „**Zasada dwóch trzecich**“ w naukach technicznych mówi, że podwojenie mocy produkcyjnych wiąże się najczęściej ze wzrostem kosztów materiałowych jedynie o dwie trzecie.
- **Stochastyczne korzyści skali** wynikają z tego, że wraz z rosnącą wielkością przedsiębiorstwa w wyniku prawa wielkich liczb łatwiej jest kalkulować wydarzenia przypadkowe. Pozwala to na utrzymywanie mniej kosztownych stanów magazynowych, względnie rezerw czynników produkcji.
- Przy różniącej się minimalnej efektywnej skali dla **następujących po sobie etapów produkcji** otrzymamy optymalny potencjał całkowity dopiero przy najmniejszej wielokrotności wartości minimalnej efektywnej skali dla poszczególnych etapów.

²⁰ Subaddytywna funkcja kosztów zdefiniowana jest w ten sposób, że suma kosztów K_n każdego oferenta n , który posiada udział σ_n w ogólnej produkcji, przekracza koszty pojedynczego oferenta K :

$$K(x) < \sum_n^N K_n(\sigma_n x) \text{ mit } \sum_n^N \sigma_n = 1; \sigma_n \geq 0.$$

²¹ Rosnące przychody ze skali zdefiniowane są tylko w przypadku stałego stosunku czynników wejściowych. Spadające koszty średnie mogą powstawać nawet wtedy, gdy stosunek czynników wejściowych zmienia się.

W przypadku wielu dóbr jakość wykonania i szybkość produkcji zależna jest od dotychczas wyprodukowanej ilości (**efekt uczenia się**), przez co w miarę upływu czasu powstają korzyści skali, ponieważ krzywa kosztów średnich przesuwana się w dół.

5.2.3 Korzyści zakresu (economies of scope)

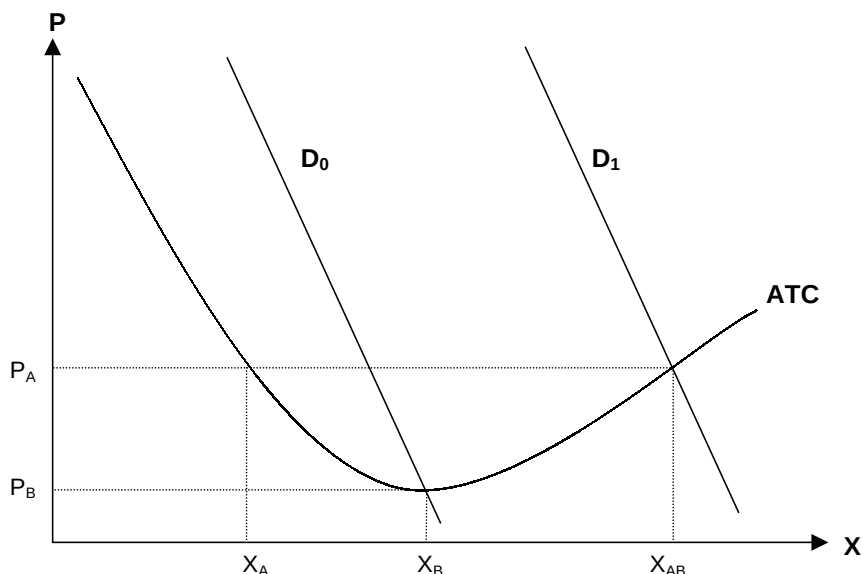
Korzyści zakresu (economies of scope) wynikają z korzystniejszej sytuacji kosztowej przy produkcji różnych dóbr. Wyrastają one z produkcji skojarzonej, niewykorzystanych mocy produkcyjnych lub efektów portfelowych. Procesy produkcji skojarzonej charakteryzują się tym, że podczas produkcji jednego dobra automatycznie produkowane są inne dobra. I tak np. podczas rafinacji ropy naftowej powstają oprócz benzyny także różne frakcje olejowe. Korzyści zakresu przy **niewykorzystanych mocach produkcyjnych** mają miejsce wtedy, gdy można je wykorzystać także do innej produkcji, jak np. sieci telekomunikacyjne (telefonia, telewizja, internet). **Efekty portfelowe** odgrywają istotną rolę, szczególnie w dziedzinie badań i rozwoju, jako wyjaśnienie niepodzielności. Przedsiębiorstwo wielobranżowe może podzielić związane z badaniami i rozwojem ryzyka techniczne i rynkowe poprzez równoczesne konstruowanie różnych produktów. Korzyści zakresu występują zawsze wtedy, gdy równoczesna produkcja wielu dóbr jest tańsza niż ich produkcja osobno.

Korzyści skali i zakresu nie muszą występować równocześnie. Korzyści zakresu – w przeciwieństwie do korzyści skali – nie prowadzą automatycznie do koncentracji na rynku po stronie oferenta. Wielu konkurentów może równocześnie utrzymywać się na rynku, jeżeli produkują pakiet dóbr charakteryzujący się korzyściami zakresu. Korzyści zakresu nie muszą powodować powstawania naturalnego monopolu. Do tego konieczne są jeszcze korzyści skali.

Poniższy wykres przedstawia sposób oddziaływania korzyści zakresu. Popyt przedstawiony jest na początku w postaci krzywej D_0 . Założono, że oferenci A i B mieliby taką samą funkcję średnich kosztów całkowitych (ATC) i ustalili ceny na podstawie swoich kosztów średnich. Równocześnie przedsiębiorstwo B ma większe zasoby niż A i odpowiednio jest w stanie zwiększyć swoje moce produkcyjne poprzez inwestycje. W początkowym okresie A może zaoferować ilość X_A po cenie P_A . Oferent B jest w stanie – w odróżnieniu od A – zaoferować całą poszukiwaną przez

konsumentów ilość X_B i to za niższą cenę P_B z powodu subaddytywności kosztów w odpowiednim (poszukiwanym) zakresie produkcji. Przedsiębiorstwo B wyruguje z rynku przedsiębiorstwo A i będzie mogło korzystać z pozycji monopolisty na niekorzyść przeciwnej strony rynku.

Wykres 21: Ilość i korzyści skali



5.2.4 Skończoność korzyści skali i zakresu

Korzyści skali i zakresu mogą się wyczerpywać, co prowadzi do tego, że koszty średnie od pewnej wielkości produkcji – X_B na Wykres 21 – przestają spadać a nawet ponownie wzrastają, czego przyczyny należy upatrywać szczególnie w kosztach transakcyjnych wewnątrz przedsiębiorstwa. Rosną one nieproporcjonalnie wraz z rozrostem przedsiębiorstwa, na przykład, dlatego że organizacja podziału pracy jest kosztowniejsza i pojawiają się nieefektywności organizacyjne.

W połączeniu z **rosnącym popytem** ograniczony charakter korzyści skali może doprowadzić do tego, że istniejący w danym momencie monopol naturalny odejdzie w przeszłość. Podczas, gdy przy wielkości popytu D_0 istnieje naturalny monopol, to przy wzroście wielkości popytu do D_1 na rynku może na stałe pozostać dwóch oferentów, ponieważ dwóch oferentów jest w stanie wytworzyć zwiększoną, łączną wielkość popytu X_{AB} po niższych kosztach niż monopolista A w fazie wyjściowej,

który mógł sprzedawać wyłącznie po cenie wyższej, czyli P_A . Podobnie wygląda sytuacja w przypadku postępu technicznego, na skutek którego minimalne koszty jednostkowe mogą być osiągnięte już przy zdecydowanie mniejszej wielkości produkcji. W ten sposób można wytłumaczyć małe zainteresowanie innowacjami naturalnego monopolisty, który w przeciwnej sytuacji ponosiłby ryzyko utraty swojej pozycji rynkowej.

5.2.5 Wnioski

Zawodność rynku związana z niepodzielnością występuje z największą siłą na rynkach, które oprócz subaddytywności charakteryzują się istotną nieodwracalnością procesów, tj. wysokimi inwestycjami, które są nieodzowne dla działalności na określonym rynku (a w przypadku odejścia z rynku stanowiłyby koszty utopione (**sunk costs**)²²). Na takich rynkach występuje małe zainteresowanie potencjalnej konkurencji wejściem na rynek. Opisana sytuacja często ma miejsce w zmonopolizowanych obszarach tzw. „wąskich gardeł”. W takich sytuacjach wskazane są z reguły ingerencje państwa celem korekty nieefektywnej alokacji i uniemożliwienia wykorzystywania siły rynkowej. Tabela 2 przedstawia niepodzielność i nieodwracalność dla sektorów, dla których często zakłada się duże prawdopodobieństwo zawodności rynku. Także na rynkach cechujących się wysoką nieodwracalnością i małą subaddytywnością może - w zależności od uwarunkowań rynkowych - dojść do zakłóceń konkurencji, podczas gdy przy wysokiej subaddytywności i małej nieodwracalności możliwe monopolie naturalne będą dostatecznie dyscyplinowane przez potencjalną konkurencję. W takiej sytuacji, inaczej niż przy wysokich kosztach utopionych, występuje duża zachęta do krótkoterminowego wchodzenia na rynek (tzw. strategia „hit and run”). Im ciemniejszym kolorem oznaczono pole tabeli, tym bardziej należy spodziewać się ingerencji państwa na rynku. W przypadku oligopolistycznej struktury rynku należy zwrócić szczególną uwagę na niezgodne z prawem konkurencji porozumienia uczestników rynku. Mają tu znaczenie nie tylko porozumienia co do cen, lecz także

²² Koszty utopione to koszty nieudanego wejścia na rynek lub – mówiąc technicznie – koszty własnej decyzji nieodwracalnych składników kosztów przy odchodzeniu z rynku. Kosztów utopionych nie można utożsamiać z kosztami stałymi, ponieważ przy odejściu z rynku część kosztów stałych będzie można sprzedać (np. nieruchomości – uznawane są one za koszty stałe, ale wraz z odejściem z rynku mogą zostać sprzedane i nie będą zaliczone do kosztów utopionych).

uzgodnienia co do identycznych metod kalkulacji cenowej, warunków sprzedaży lub rabatów.

Tabela 2: Subaddytywność i nieodwracalność w różnych dziedzinach gospodarki

Sektor	Szczegół produkcji	Subaddytywność	Nieodwracalność
Energia elektryczna/Gaz	Wytwarzanie	Nie	Umiarkowana
	Dystrybucja	Tak	Wysoka
Ciepło przemysłowe	Produkcja	Nie	Wątpliwa
	Dystrybucja	Tak	Wysoka
Woda	Produkcja	Nie	Niewielka
	Dystrybucja	Tak	Wysoka
Odpady	Zbieranie	Tak	Mała
	Spalanie	Wątpliwa	Wysoka
Telewizja kablowa	Program	Nie	Mała
	Dystrybucja	Tak	Wysoka
Listy, paczki	Transport	Wątpliwa	Wątpliwa
	Dostarczanie	Tak	Mała
Koleje	Sieć	Tak	Wysoka
	Przewozy	Nie	Mała
Linie autobusowe, samochody ciężarowe		Nie	Mała
Żegluga		Mała	Mała
Rurociągi		Tak	Wysoka

5.3 Ograniczenia informacyjne

Akerlof, G.A. (1970): The Market for "Lemons" – Quality Uncertainty and the Market Mechanism. In: Quarterly Journal of Economics, vol. 84, str. 488 – 500. Alchian, A. A. (1970): Information Costs, Pricing, and Resource Unemployment. In: Phelps, S. et al. (Eds.): Microeconomic Foundations of Employment and Inflation Theory. New York, str. 27-52. Arrow, K. J. (1963): Uncertainty and Medical Care. In: American Economic Review, Vol. 53, str. 941-973. Farrell, J.; Klemperer, P. (2005): Coordination and Lock-In: Competition with Switching Costs and Network Effects. In: Handbook of Industrial Organization, Vol. 3. Amsterdam. Shapiro, C. (1983): Premiums for High Quality Products as Returns to Reputation. In: The Quarterly Journal of Economics, Vol. 98, str. 659-679.

W modelu konkurencji doskonałej przyjmuje się założenie, że wszyscy uczestnicy informowani są wyczerpująco, na czas i nieodpłatnie. W praktyce jednak uczestnicy rynku są często niedoinformowani w takim stopniu, że zagraża to poważnie funkcjonowaniu rynku i może doprowadzić do jego zawodności. Zagrożenie zawodnością rynku zależy w dużej mierze od charakteru **dóbr** będących przedmiotem wymiany rynkowej. W przypadku towarów bazujących na doświadczeniu i zaufaniu²³ należy liczyć się z ograniczeniem właściwego funkcjonowania rynku, a z kolei w przypadku towarów jednorodnych nie należy się spodziewać zakłóceń w wyniku ograniczeń informacyjnych.

Generalnie w ramach ograniczeń informacyjnych należy rozróżniać pomiędzy niewiedzą a niepewnością.

5.3.1 Niewiedza

Nieznajomość korzyści występuje najczęściej w odniesieniu do dóbr niematerialnych (choćby ochrona zdrowia lub edukacja), których korzyści ujawniają się dopiero po długim czasie. Koszty podaży takich dóbr są znane, nabywcy nie potrafią jednak oszacować korzyści i rezygnują z konsumpcji. Rynek zawodzi, w związku z czym na rynek ten w roli oferenta wkracza państwo.

Nieznajomość ceny odnosi się do racjonalnej nieznajomości ceny równowagi, poprzez co transakcje są opóźniane lub zawierane są po cenach, które nie prowadzą do opróżnienia rynku. Cena „czyszcząca” rynek ustala się z reguły sukcesywnie, przy

²³ W przypadku towarów bazujących na doświadczeniu ich jakość jest w pełni znana dopiero po ich konsumpcji, gdyż zbadanie jej przed zawarciem transakcji możliwe jest przy poniesieniu znacznych kosztów (produkty spożywcze). W przypadku towarów bazujących na zaufaniu pełne zbadanie jakości nie jest możliwe ani przed ani po ich konsumpcji (lekarstwa).

czym najpierw dochodzi do wielu transakcji po cenach suboptymalnych. Czas trwania **procesów dopasowania** uzależniony jest istotnie od stanu poinformowania jednostek oraz nakładów, jakie muszą poczynić w celu uzyskania odpowiedniej informacji. Ponieważ zdobywanie informacji wiąże się z kosztami, z optymalną sytuacją mamy do czynienia wówczas, gdy cena pozyskania kolejnej jednostki informacji odpowiada oczekiwanym oszczędnościom. W takich przypadkach nie dochodzi do zawodności rynku. Inaczej jest jednak, gdy w następstwie niepodzielności lub efektów zewnętrznych powstają trudności w dostępie do informacji. Mogą wtedy pomóc systemy informacji o rynku, które zapobiegają zawodności rynku.

5.3.2 Niepewność, asymetria informacyjna

Wątpliwe jest, by rynek mógł zawieść z powodu **niepewności** przedsiębiorców, ponieważ istnieje cały szereg możliwości zabezpieczenia się przed następstwami nieoczekiwanych zachowań poprzez odpowiednio skonstruowane umowy lub zawarcie stosownego ubezpieczenia. Inaczej dzieje się w przypadku **asymetrycznie rozdzielonej informacji**²⁴, w którym to przypadku może dojść do opisanych poniżej rodzajów zawodności rynku:

Jeżeli uczestnicy rynku są źle poinformowani o okolicznościach istotnych dla dokonania transakcji, negatywna selekcja (**adverse selection**) prowadzi do załamania rynku dóbr wysokiej jakości i pozostania na nim wyłącznie dóbr o jakości niskiej, ponieważ kiedy odbiorcy znają jedynie przeciętną jakość oferowanych dóbr, nie są gotowi do zapłacenia wyższej ceny za domniemane dobro wysokiej jakości. W następstwie zanikają zachęty do oferowania dóbr wyższej jakości, dominująca na rynku przeciętna jakość stale się obniża, a co za tym idzie, spada także gotowość do zapłaty. Taki proces prowadzi do ww. wyniku, tzn. załamuje się rynek towarów

²⁴ **Teoria agencji** przedstawia powszechnie stosowane ramy analizy problemów związanych z asymetrycznie rozdzielonymi informacjami. Podstawową cechą jest tu odróżnienie pryncypała, zleceniodawcy i agenta, który otrzymuje zlecenie podjęcia określonych działań. Przedstawione są tam warunki, które muszą być spełnione, aby w sytuacji nierównego dostępu do informacji oczekiwana transakcja mogła dojść do skutku na tle utajnionych działań wzgl. informacji („hidden action” wzgl. „hidden information”), utajnionych właściwości („hidden characteristics”) oraz utajnionych zamiarów („hidden intentions”). Podkreśla się przy tym przede wszystkim znaczenie przypadkowych wpływów.

wysokiej jakości (w literaturze ekonomicznej karierę zrobił przykład rynku samochodów używanych).

Problem pokusy nadużycia (**moral hazard**) względnie oportunistycznego poumownego może wystąpić, jeśli po zawarciu kontaktu nie ma możliwości dokładnej kontroli tego, w jakim stopniu strony kontraktu wywiązały się z umowy. Ponieważ nie wszyscy nabywcy zachowują się zgodnie z kontraktem względnie z należytą starannością, a oferenci kalkulują swoje ceny w taki sposób, aby nie ponieść straty, w cenach pojawia się premia za ryzyko. W ten sposób uczestnicy spełniający warunki kontraktu obciążeni są względnie wysoką ceną i dlatego skłonni są rezygnować z nabywania pewnych dóbr, a w ekstremalnej sytuacji nie są zawierane dalsze transakcje z uczciwymi uczestnikami rynku i rynek zawodzi, ponieważ z powodu wysokiej ceny na rynku pozostają jedynie oportuniści. To odroczone zagrożenie odnosi się przede wszystkim do sposobu funkcjonowania rynków ubezpieczeń, obciążonych takim poumownym oportunistycznym w formie „złych” ryzyk.

Dalsze niedoskonałości rynkowe spowodowane ograniczeniami informacyjnymi mogą wystąpić (najczęściej przy kontraktach długoterminowych) w wyniku stosunku kontraktowego, gdy umowa zawiera nieścisłości interpretacyjne i przynajmniej jedna ze stron kontraktu poniosła specyficzne, nieodwracalne inwestycje. Wtedy występuje niebezpieczeństwo, że jedna ze stron kontaktu będzie interpretowała umowę jednostronnie na swoją korzyść i wstrzyma się z wykonaniem części świadczeń lub zażąda większego świadczenia od drugiej strony kontraktu, co w literaturze określa się mianem „**postępowania hold up**”. W przypadku zerwania stosunku umownego powstałyby koszty utopione – „sunk cost”, co uzależnia, a tym samym umożliwia wyzysk danego uczestnika rynku. Takie zależności spowodowane nieodwracalnymi inwestycjami nazywane są **efektem lock in** („efektem pochwycenia”). Inaczej niż w przypadku pokusy nadużycia niepożądane zachowania są tu łatwe do wykrycia. Brakuje jednak odpowiednich sankcji, które nakłoniłyby partnera kontraktowego do lojalnego zachowania. Jeżeli prawdopodobieństwo wystąpienia oportunistycznego zachowania partnera kontraktu nie może zostać dobrze oszacowane przy zawieraniu kontraktu, to nie powinny dochodzić do skutku określone inwestycje względnie transakcje. Dostęp osób trzecich do rynku będzie zatem utrudniony w sytuacji występowania efektu „pochwycenia”.

Jeżeli ograniczenie informacyjne polega na nierówno rozłożonej względnie asymetrycznej informacji, to nie musi to koniecznie prowadzić do zawodności rynku, ponieważ dla kupujących pozostaje możliwość doinformowania się, (tzw. **screening**), a sprzedający mogą uwiarygodniać swoją prawdziwą jakość za pomocą odpowiednich sygnałów (tzw. **signaling**). W dalszej kolejności obie strony rynku mogą próbować uwiarygodnić własną jakość poprzez osiągnięcie dobrej reputacji, albo w drodze zawierania umów zawierających zapisy o powstrzymywaniu się od określonych szkodliwych praktyk, albo przewidujących rabaty, jeśli produkt okaże się pozbawiony wad. Również powiązania wertykalne umożliwiają uniknięcie zawodności rynku będącej konsekwencją asymetrii informacyjnej. Kontrola fuzji powinna w takich przypadkach zwrócić uwagę na niedoskonałości rynku wiążące się z przepływem informacji.

B. Zastosowanie ekonomii w prawie konkurencji

6 Wyznaczanie rynku

6.1 Wprowadzenie

Wyznaczenie rynku właściwego stanowi centralny element dokonywanej na podstawie prawa konkurencji oceny koncentracji przedsiębiorstw, przypadków nadużywania pozycji dominującej oraz pionowych i poziomych ograniczeń konkurencji. Prawidłowa ocena siły rynkowej w przypadkach zachowań o charakterze nadużyć albo w przypadkach łączenia przedsiębiorstw wymaga, by konieczność ewentualnej interwencji została właściwie oceniona. Punktem wyjścia jest produkt lub usługa, oferowane przez brane pod uwagę przedsiębiorstwo. Dla osądu istniejących stosunków konkurencji, produkty lub usługi muszą zostać zidentyfikowane z punktu widzenia ich substytucyjności.²⁵ Zadanie wyznaczenia rynku polega tym samym na systematycznym badaniu sił konkurencji, którym podlegają brane pod uwagę przedsiębiorstwa. Wynik wyznaczenia rynku może wypaść różnie w zależności od sformułowania zapytania (np. odnośnie fuzji czy porozumień wertykalnych).

Wyznaczenie rynku ma każdorazowo wymiar przedmiotowy i geograficzny. Wedle stosowanej przez Komisję Europejską definicji, rynek właściwy produktowo obejmuje wszelkie wyroby i/lub usługi, które w ocenie konsumentów są zamienne lub substytucyjne z punktu widzenia ich właściwości, cen oraz zastosowania. Rynek właściwy geograficznie obejmuje obszar, na którym zaangażowane przedsiębiorstwa oferują właściwe produkty lub usługi, na którym warunki konkurencji są wystarczająco jednolite i który od obszarów sąsiadujących odróżnia się odczuwalnie odmiennymi warunkami konkurencji.²⁶

²⁵ Dotyczy to również podaźowej strony rynku. Wypracowane na użytek analizy podaźowej siły rynkowej zasady, dotyczące określenia rynku właściwego w zasadzie można symetrycznie zastosować do strony popytowej rynku (określenie rynków nabywczych perspektywy odbiorcy).

²⁶ Por. European Commission, Commission notice on the definition of relevant market for the purpose of Community competition law (97/C 273/03), para. 7,8 (see at [http://eur-lex.europa.eu/LexUriServ/LexUriServ.do?uri=CELEX:31997Y1209\(01\):EN:HTML](http://eur-lex.europa.eu/LexUriServ/LexUriServ.do?uri=CELEX:31997Y1209(01):EN:HTML)).

Wyznaczenie rynku nie może być traktowane jako cel sam w sobie. Jest to raczej instrument pozwalający określić relacje konkurencyjne i ich granice. Wraz z definicją rynku właściwego powstają ramy analityczne, będące bazą dla oceny konkurencji i ułatwiające tę ocenę. W niektórych przypadkach wyznaczenie rynku ma decydujący wpływ na ocenę przypadku. Warunkiem udowodnienia pozycji dominującej jest wyznaczenie rynku właściwego, które umożliwia obliczenie udziału w rynku, ważnego wskaźnika siły rynkowej. Na tym polega znaczenie określania rynku właściwego.

Kim są konkurenci branych pod uwagę przedsiębiorstw? W wielu przypadkach odpowiedź na to pytanie, a co za tym idzie także prawidłowe wyznaczenie rynku właściwego, nie są sprawą oczywistą. Czy np. przedsiębiorstwo kolejowe, oferujące usługi transportu węgla, konkuruje na rynku kolejowego transportu węglowego, na rynku kolejowego transportu masowych dóbr sypkich, usług transportowych dóbr sypkich koleją, transportem drogowym czy wodnym, czy na rynku innej grupy usług transportowych? W każdym przypadku istotne jest, aby dokładnie przeanalizować produkty i usługi po to, by zrozumieć sposób funkcjonowania rynków.

W konkretnych przypadkach pytanie o zakres rynku właściwego nie zawsze musi doczekać się pełnej odpowiedzi. Przypadek taki zachodzi wtedy, gdy ocena w aspekcie przepisów prawa konkurencji dla każdego możliwego rynku właściwego pozwala stwierdzić, iż nie powstaje zagrożenie dla konkurencji. W takiej sytuacji kwestia wyboru pomiędzy alternatywnymi definicjami rynku A i B może nierzadko pozostać otwarta.

Wyznaczanie rynku właściwego produktowo i geograficznie ułatwia szereg metod bądź koncepcji. W następnym punkcie omówione zostanie bliżej wyznaczanie rynku właściwego produktowo.

6.2 Metody i koncepcje jakościowe

Dwie główne, podstawowe koncepcje używane przy wyznaczaniu rynku właściwego to substytucyjność popytu oraz substytucyjność podaży. Z ekonomicznego punktu widzenia wszelkie produkty i usługi konkurują o siłę nabywczą konsumentów. Z perspektywy konsumenta istnieje większa substytucyjność między jednymi produktami (np. lemoniada i woda mineralna) niż ma to miejsce w przypadku innych

wyrobów (np. lemoniada i oprogramowanie komputerowe). Innymi słowy, łańcuch substytucji, w którym znajdują się wszystkie produkty, wykazuje luki substytucyjne o różnej szerokości. W praktyce należy więc wyjść z założenia, że, po pierwsze, konkurujące w ramach jednego rynku produkty nie są stuprocentowymi substytutami, z drugiej zaś, że przy produktach nie należących do właściwego rynku produktowego, które jednak dają się zaliczyć do obszaru bliskiego rynkowi, substytucja nie jest zerowa. Szczególnie w przypadku zróżnicowanych produktów i istniejących preferencji konsumentów stopniowe różnice substytucyjności mogą mieć duże znaczenie. Precyzyjne wyznaczenie rynku właściwego produktowo może zatem okazać się trudne.

Szczególnie w przypadkach, w których ściśle wyznaczenie produktowego rynku właściwego ma duże znaczenie, zasadnym jest, by w ramach oceny konkurencyjności uwzględnić również tzw. obszary bliskie rynkowi, ewentualnie otoczenie konkurencyjne (por. r. 7.1.6). Produkty z obszaru bliskiego rynkowi nie należą do rynku właściwego. Zalicza się do niego natomiast takie produkty, które są substytucyjne w długim okresie z punktu widzenia konsumenta, w stosunku do produktów zaliczanych do rynku właściwego. Produktami z obszaru bliskiego rynkowi mogą być ponadto takie dobra, które z punktu widzenia oferenta, co najmniej w średnim okresie, charakteryzują się wysoką substytucyjnością podażową.

Wysoki stopień substytucyjności pomiędzy dwoma produktami jest prawdopodobny wtedy, gdy ich obiektywne (np. techniczne) właściwości i cel zastosowania są bardzo podobne. Wskazuje to na funkcjonalną lub ewentualnie gospodarczą²⁷ zamiennieść (i vice versa). Przy ocenie substytucyjności pomiędzy produktami należy zwrócić uwagę na fakt, że nawet gdy obiektywne właściwości i cel użytkowy produktów są do siebie podobne, substytucja może być ograniczona²⁸. Przeszkodę dla zamiennieści mogą stanowić np. wierność marce w odniesieniu do określonych towarów rynkowych, istniejące koszty zmiany (switching costs) lub ograniczenia prawne. Przykładowo, oparta na preferencjach wierność marce drogich perfum może

²⁷ Gospodarcza wymiennieść wymaga dodatkowo podobieństwa cen.

²⁸ Por. ICN Merger Working Group: Investigation and Analysis Subgroup, ICN Merger Guidelines Workbook, April 2006, A. Market Definition, p. 20 ff.

usprawiedliwiać określenie odrębnego rynku dla tych perfum²⁹. Także, kiedy koszty zmiany (np. koszty przebrojenia, wymagania inwestycyjne) są wysokie w relacji do ceny produktu, substytucyjność staje się mniej prawdopodobna. Np. w dziedzinie handlu detalicznego przedsiębiorstwa handlowe usiłują, poprzez tzw. karty lojalnościowe wpłynąć na koszty zmiany swoich klientów, tj. podnieść je.

Ważne informacje o istniejących stosunkach substytucji zawierają (względne) ceny oraz ich zmiany. Na przykład, równoległy rozwój cen dwóch produktów w czasie, niewynikający z analogicznego kształtowania się kosztów ich wytworzenia, jest pierwszą wskazówką, że mogą być to bliskie substytuty. Substytucyjność popytową pomiędzy dwoma lub więcej produktami można analizować za pomocą tzw. testu hipotetycznego monopolisty, w którym cena odgrywa centralną rolę (por. r. 6.3.1).

Przy zastosowaniu koncepcji substytucyjności popytu odwołuje się często do „**przeciętnego**”, ewentualnie „**rozsądnego**” konsumenta. Nie można jednak zapomnieć, że najczęściej klienci nie są identyczni, lecz zróżnicowani. Ekonomisci używają pojęcia **konsumenta krańcowego** (marginal customer), który przy podwyżce cen opuszcza rynek jako pierwszy. Nawet, gdy w niektórych przypadkach przeciętny konsument jest „uwięzionym” konsumentem (captive customer), istnieją potencjalni konsumenci graniczni, którzy są gotowi przy relatywnej zmianie cen wybrać substytut. Przy wyznaczaniu rynku chodzi ostatecznie o pytanie, czy jest wystarczająco dużo konsumentów granicznych gotowych wybrać produkt substytucyjny tak, aby podwyżka cen stała się dla oferentów nieopłacalna. W niektórych przypadkach, na przykład w regularnym lotniczym ruchu pasażerskim, istnienie rozmaitych typów klientów ewentualnie grup (np. pasażerowie wrażliwi na cenę *kontra* wrażliwi na czas podróży) może prowadzić do wyodrębnienia dwóch rynków produktowych.³⁰

²⁹ Wskazówki odnośnie preferencji klientów i istniejące wzory konsumpcji można uzyskać poprzez ankiety wśród klientów lub z dostępnych badań zachowań konsumenckich, prowadzonych w instytutach badań rynkowych.

³⁰ European Commission, 11.02.2004, Case COMP/M.3280 - "Air France/KLM", para. 19.

Koncepcję substytucyjności popytu uzupełnia **koncepcja substytucyjności podaży** („substytucyjność/zamienność” także po stronie oferentów)³¹. Substytucyjność popytu z reguły wywiera najbardziej bezpośredni, dyscyplinujący wpływ na konkurencję na rynku. Jeżeli jednak oferenci są w stanie, nawet przy małych, trwałych zmianach cen, przestawić swoją produkcję na dane wyroby i w krótkim terminie wprowadzić je na rynek, wtedy również elastyczność podaży działa w sposób dyscyplinujący na zachowanie uczestników rynku. Zdarza się to często wtedy, kiedy oferowane są różnorodne rodzaje bądź typy jednego produktu, które nie są zamienne dla konsumenta³². Można sobie więc wyobrazić, że wytwórcy określonych rodzajów i wymiarów śrub są w stanie przestawić produkcję na inne rodzaje i wymiary. Kiedy takie przestawienie po stronie oferenta jest szybkie i bezproblemowe, zachodzi substytucyjność podaży, w związku z czym wszystkie rodzaje i wymiary śrub mogą zostać zaliczone do tego samego rynku właściwego.

W określonych przypadkach, gdy rozpatrywane towary bądź usługi były już przedmiotem wcześniejszych postępowań organu ochrony konkurencji, istnieje możliwość sięgnięcia po wyniki dokonanego wcześniej określenia rynku właściwego. Ma to miejsce wtedy, gdy wykształciła się stała praktyka administracyjna względem wyznaczania danego rynku, jak też gdy owa praktyka została potwierdzona przez właściwe sądy. Gdy taki przypadek nie zachodzi, pomocnym może być powołanie się na już istniejące postępowania innych organów ochrony konkurencji (można przykładowo przyjąć wypracowaną przez nie definicję rynku jako hipotezę roboczą). Jednak przed przejęciem wyników innych postępowań należy sprawdzić, czy określenie rynku może być rzeczywiście przeniesione do badanego aktualnie przypadku. Takie postępowanie może być niewłaściwe choćby wtedy, gdy z upływem czasu zmieniły się stosunki konkurencji, lub gdy na różnych obszarach geograficznych istnieją odmienne uwarunkowania konkurencyjne bądź regulacyjne³³.

³¹ Ibid., para. 20 ff.

³² Bundeskartellamt, powołując się na koncepcję substytucyjności podaży, wyznaczył np. jednolity produktowy rynek właściwy syntetycznie wytwarzanej witaminy C w czystej formie – pomimo istniejących różnic między produktami (patrz. Bundeskartellamt, 31.07.2002, B 3 – 27702 – "BASF/NEPG", para. 28, 36).

³³ Por. ICN Merger Working Group: Investigation and Analysis Subgroup, ICN Merger Guidelines Workbook, April 2006, A. Market Definition, p. 18.

6.3 Metody ilościowe – test SSNIP

Oprócz jakościowych metod i koncepcji określania rynku właściwego zastosować można również metody ilościowe. Podczas gdy na podstawie metod jakościowych ocenia się substytucyjność na bazie charakterystycznych cech poszczególnych produktów z punktu widzenia reprezentatywnego, ewentualnie racjonalnego konsumenta, dochodząc tym samym do wniosków dyskrecjonalnych („produkt jest substytucyjny bądź nie jest substytucyjny”), to metody ilościowe umożliwiają dodatkowo ocenę stopnia substytucyjności produktów. W przypadku użycia metod ilościowych, analiza polega nie tylko na ocenie szczególnych cech produktów oraz hipotez na temat substytucyjności z punktu widzenia konsumenta, lecz ocenia owe hipotezy na bazie rzeczywistych danych rynkowych. Często daje się przez to uzyskać precyzyjniejszą ocenę zamienności produktów. Metody ilościowe stanowią więc uzupełnienie dla jakościowych badań rynku. Poza tym taka forma badań opiera się na rzeczywistych (agregowanych) zachowaniach odbiorców, a co za tym idzie w wynikach uwzględnione zostają aspekty, które przy analizie funkcji produktów mogłyby pozostać nierozpoznane.

6.3.1 Koncepcja podstawowa

Koncepcją ilościową służącą do wyznaczania rynku jest stosowany przez wiele organów ochrony konkurencji tzw. **test SSNIP**, zwany też **testem hipotetycznego monopolisty**. Celem testu SSNIP jest określenie rynku właściwego. Test SSNIP może zostać przyjęty zarówno do wyznaczania rynku produktowego, jak też rynku geograficznego.

Podstawowym zagadnieniem testu SSNIP jest kwestia, czy (hipotetycznie monopolistyczny) oferent mógłby zwiększyć swój zysk poprzez trwały wzrost ceny danych produktów o średnio 5-10% (SSNIP = Small Significant Non-transitory Increase in Price)³⁴. Gdyby to okazało się dlań zyskowne, oznacza to, że klienci przy wzroście cen tylko w niewielkim stopniu przesuną się w kierunku produktów

³⁴ Założenie istnienia tylko jednego producenta na rynku jest ważne, gdyż pomaga poznać działanie konkurencji w ramach rozpatrywanej grupy produktów. Alternatywnie można by też przyjąć, że wszyscy oferenci danego produktu wspólnie podnoszą ceny.

alternatywnych. Owe produkty stanowią zatem odrębny rynek produktowy. Jeżeli natomiast konsumenci w istotnym stopniu skłaniają się ku produktowi alternatywnemu, uznając go za zamienny, to taka podwyżka byłaby dla monopolisty nieopłacalna. W takim przypadku substytucyjność wobec innego produktu byłaby wysoka, i należałoby produkt taki zaliczyć do wyznaczanego produktowego rynku właściwego. Kolejnym etapem jest sprawdzenie tą samą metodą, czy jeszcze inne produkty należą do danego rynku właściwego. Wówczas znowu zakłada się, że produkty na rynku oferowane są przez (hipotetycznego) monopolistę. Gdyby podwyższenie ceny własnych produktów o 5-10%, również było dla niego nieopłacalne z uwagi na zwiększone zainteresowanie konsumentów produktem alternatywnym, to również ów alternatywny produkt zaliczany jest do produktowego rynku właściwego. Produktowy rynek właściwy zostaje wyznaczony wtedy, gdy wzrost cen staje się dla monopolisty opłacalny; w takim wypadku żaden produkt alternatywny nie stanowi substytutu dla wystarczająco dużej liczby konsumentów. Punktem wyjścia jest tym samym wyznaczenie rynku właściwego w sposób maksymalnie zawężony. Właściwy rynek produktowy jest następnie krok po kroku poszerzany do momentu, kiedy SSNIP przestaje być opłacalny dla (hipotetycznego) monopolisty³⁵.

Sposób postępowania stosowany do ustalenia właściwego rynku produktowego można przenieść na grunt wyznaczania geograficznego rynku właściwego. Jeśli przy podwyżce rzędu 5-10% znaczna część konsumentów (strony popytowej rynku) wybiera produkty z sąsiadującego regionu, podwyżka taka nie byłaby opłacalna dla hipotetycznego monopolisty. W takim przypadku ów sąsiadujący region należałoby również zaliczyć do rynku właściwego. Gdyby natomiast podwyżka była opłacalna, a co za tym idzie popyt nie przeniósłby się do sąsiadujących regionów, to regiony te nie należałyby do wyznaczanego rynku właściwego.

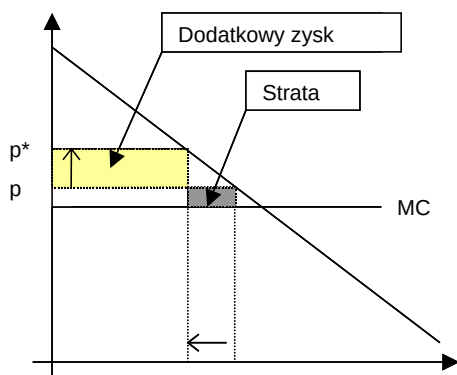
Powyższe przykłady rozpatrują **zamiennność produktów** z punktu widzenia **nabywcy**. Jednak test SSNIP obejmuje nie tylko zamiennność produktów z konsumenckiego punktu widzenia, ale bierze też pod uwagę **substytucyjność**

³⁵ Owo stopniowanie procedury kryje w sobie niebezpieczeństwo, iż wynik testu uzależniony będzie od kolejności branych pod uwagę produktów. Stąd też należy zwracać uwagę na testowanie rozmaitych kombinacji.

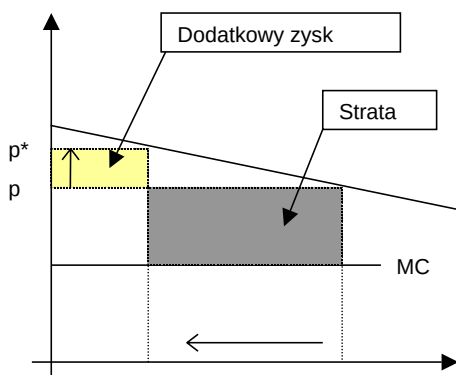
podaży. Jeżeli istnieje przedsiębiorstwo wytwarzające produkt, który wprowadzić nie jest zamienny z produktem monopolisty, ale jego wytwórca dysponuje odpowiednimi technologiami, prawami i wiedzą, aby dany produkt w każdej chwili stworzyć, to oferent ów również wywiera nacisk konkurencyjny na monopolistę. Gdyby monopolista podniósł cenę swoich produktów o 5-10%, to ów oferent zdecydowałby się przestawić swoją produkcję i również zaoferować odpowiedni produkt tak, że podwyżka cen stałaby się nieopłacalna. Rynek właściwy obejmuje zatem również tego oferenta wraz z jego produktem. Gdyby natomiast ów inny oferent przy podwyżce cen nie zdecydował się na przestawienie swojej produkcji, to substytucyjność podaży można uznać za znikomą; oferent ze swoim produktem nie zalicza się do rynku właściwego.

Liczba klientów, która przy 5-10% wzroście cen odchodzi do produktów alternatywnych, zależy przede wszystkim od **elastyczności cenowej popytu**. Gdy elastyczność cenowa jest wysoka, to przy wzroście cen oferent traci stosunkowo wysoką liczbę klientów, przy niskiej elastyczności cenowej natomiast tylko niewielką ich liczbę. Jeżeli elastyczność cenowa jest większa niż 1, to wzrost ceny o 5% prowadzi do utraty więcej niż 5% klientów, a więc do zmniejszenia przychodów. To, czy wzrost ceny w tym przypadku rzeczywiście będzie korzystny dla hipotetycznego monopolisty, zależy nie tylko od elastyczności cenowej, ale też od wysokości jego **marży zysku** przed podwyżką (różnicy między ceną a kosztami zmiennymi). Pokazuje to poniższe wykresy. MC oznacza na nich koszty krańcowe lub koszty zmienne oferentów, „p” cenę przed SSNIP, a „p*” cenę po SSNIP. Na rys. 1 popyt jest dość nieelastyczny a marża zysku (różnica między „p” i „MC”) niewielka. Jak widać, dodatkowy zysk ze SSNIP (jasny obszar) wyraźnie przeważa nad stratą wynikającą z redukcji liczby sprzedanych sztuk produktu (ciemny obszar). SSNIP byłby tym samym opłacalny dla hipotetycznego monopolisty. Na rys. 2 sprawa przedstawia się odwrotnie. Popyt jest bardzo elastyczny a marża zysku wysoka. SSNIP wiąże się z wysokim spadkiem liczby sprzedanych sztuk produktu. Dodatkowy zysk ze SSNIP, porównany ze stratą związaną z mniejszą liczbą sprzedanych jednostek produktu, jest niewielki, w związku z czym SSNIP nie jest opłacalny.

Wykres 22: Zyskowość SSNIP



Rys. 1



Rys. 2

Tym samym można stwierdzić, że SSNIP tym bardziej zwiększa zysk, im mniejsza jest elastyczność cenowa i im mniejsza jest marża hipotetycznego monopolisty. SSNIP nie prowadzi natomiast do wzrostu zysku, gdy elastyczność cenowa i marża przedsiębiorstwa są wysokie.

6.3.2 Cellophane Fallacy („błąd celofanowy”)

Zastosowanie testu SSNIP może niejednokrotnie prowadzić do **błędnych wyników**, wtedy mianowicie, gdy dane produkty oferowane są przez **przedsiębiorstwa mające silną pozycję na rynku**. Gdy zastosuje się tu test SSNIP, z reguły okaże się, że nawet niewielkie podniesienie cen nie jest dla oferenta opłacalne. Z powyższego można by wyciągnąć błędny wniosek, iż istnieją substytuty, które również należałoby zaliczyć do danego rynku. W rzeczywistości jednak oferent nie podniesie swojej ceny, ponieważ cena aktualna sytuuje się już blisko ceny monopolistycznej. Podniesienie ceny ponad to nie byłoby dla monopolisty korzystne. Fenomen ten znany jest w piśmiennictwie przedmiotu jako **Cellophane Fallacy („błąd celofanowy”)**. Test SSNIP nie nadaje się zatem do wyznaczania rynku w przypadkach związanych z nadużywaniem pozycji dominującej, gdyż prowadzi on wtedy zazwyczaj do zbyt rozległego wyznaczenia rynku właściwego. Dla monopolisty

nie będzie opłacalne podnoszenie swojej ceny o 5-10%. Tak więc test SSNIP prowadzi tylko wtedy do właściwych wyników, gdy bieżące ceny rynkowe odpowiadają cenom konkurencyjnym.

6.3.3 Wykorzystanie testu SSNIP

6.3.3.1 BEZPOŚREDNIE ZAPYTANIE

Najprostszym sposobem zastosowania testu SSNIP jest bezpośrednie zapytanie klientów o ich gotowość przejścia na produkty substytucyjne. Do tego celu służą przykładowo następujące pytania:

Czy przeszliby Państwo na stosowanie stalowych puszek do napojów (i vice-versa) gdyby cena puszek aluminiowych wzrosła o:

- 1) 5% A. tak, B. nie.
- 2) 10% A. tak, B. nie.?

Czy w przypadku podwyższenia ceny o:

- 1) 5%
- 2) 10%

na puszki do napojów oferowane przez waszego (albo najbliższego) dostawcę, zmieniliby Państwo tego dostawcę na dostawcę z siedzibą w innym miejscu?

- 1) A. tak, B. nie.
- 2) A. tak, B. nie.

Wynik ankiety daje pojęcie o tym, jak duży jest udział klientów, którzy przy wzroście ceny grupy produktów przechodzą na produkty alternatywne. Gdy udział klientów, którzy przy 5% wzroście cen zmieniają produkt jest wysoki, to podniesienie cen byłoby z wysokim prawdopodobieństwem nieopłacalne. W tym przypadku należy wyjść z założenia, że produkt alternatywny przynależy do rynku. Jeśli natomiast przy 10% wzroście cen z produktu rezygnuje jedynie niewielka część klientów, to taki wzrost cen byłby zapewne dla oferenta opłacalny, a rozpatrywany substytut w

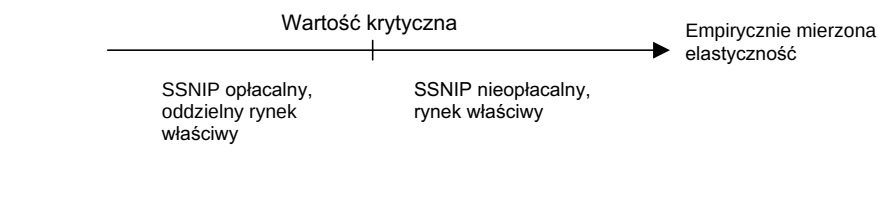
związku z tym nie byłby zaliczany do rynku właściwego. Metoda ta nie bierze natomiast pod uwagę marży realizowanej na danych produktach. Obok ankietowania klientów lub konkurentów można też wykorzystać metody ilościowe by sprawdzić, czy SSNIP jest opłacalny. Poniżej przedstawione zostaną elastyczność cenowa, elastyczność mieszana (elastyczność krzyżowa) oraz analiza korelacji cen³⁶.

6.3.3.2 ELASTYCZNOŚĆ CENOWA

Za pomocą elastyczności cenowej uzyskuje się informację o tym, jak konsumenci reagują na podwyższenie ceny produktu. Gdy elastyczność cenowa popytu wynosi 0,3, należy wnioskować, że przy 5-10% wzroście cen popyt zmaleje tylko o ok. 1,5-3%. W obliczu stosunkowo niskiego spadku popytu taki wzrost cen wydawałby się opłacalny dla hipotetycznego monopolisty; substytucyjność wobec innych produktów jest ograniczona, a tym samym rynek właściwy nie może być szerszy.

Jak wykazano wcześniej, nie tylko od elastyczności ale też od marży oferenta zależy, czy SSNIP jest korzystny dla hipotetycznego monopolisty. Im mniejsza elastyczność i marża, tym bardziej jest to opłacalne. Z informacji o aktualnej cenie i kosztach zmiennych hipotetycznego monopolisty można obliczyć tzw. „wartość krytyczną” elastyczności cenowej, przy której SSNIP akurat jeszcze jest opłacalny. W celu wyznaczenia rynku właściwego wartość tę należy porównać z uzyskaną empirycznie wartością elastyczności cenowej. Jeżeli wartość empiryczna jest mniejsza od wartości krytycznej, to podwyżka cen byłaby opłacalna, a odpowiednie produkty tym samym przynależne do oddzielnego rynku. Gdy wartość empiryczna jest większa od wartości krytycznej, to podwyżka nie byłaby opłacalna, a rynek trzeba określić szerzej. Powyższe obrazuje następujący wykres.

Wykres 23: SSNIP i elastyczność cenowa popytu



³⁶ Przykładową sprawę, w której Komisja zastosowała metody ilościowe znaleźć można w Comp/M.2187 CVS / Lenzing

6.3.3.3 ELASTYCZNOŚĆ MIESZANA (ELASTYCZNOŚĆ KRZYŻOWA) POPYTU

Elastyczność mieszana popytu wskazuje jak ściśle są relacje substytucyjności pomiędzy produktami. Umożliwia także ocenę, w jakiej mierze popyt na produkt A zmienia się, gdy zmienia się cena alternatywnego produktu B. Gdy elastyczność mieszana popytu jest wysoka, SSNIP często nie będzie opłacalny dla hipotetycznego monopolisty, ponieważ znacząca część klientów przeszłaby do oferenta produktu alternatywnego. Jest to wskazówka na to, że produkty te należą do tego samego rynku. Jeżeli natomiast elastyczność mieszana pomiędzy dwoma produktami jest niska, to należy wyjść z założenia, że produkty te należy zaliczyć do odrębnych rynków. Komisja Europejska wychodzi z założenia, że produkty, których elastyczność mieszana jest niższa niż 1, nie powinny być rozpatrywane jako substytuty.

6.3.3.4 ANALIZA KORELACJI CEN

Za pomocą analizy korelacji cen można wywnioskować, jak ściśle powiązane są zmiany cen produktów. Im większa korelacja ceny dwóch produktów, tym większe prawdopodobieństwo, że przynależą do tego samego rynku. Wstępnej oceny korelacji cen można dokonać analizując wykres zmian cen danych produktów. Gdy ceny przebiegają równolegle, wskazuje to na istnienie pomiędzy nimi ścisłego związku. Korelację pomiędzy cenami dwóch produktów mierzy się statystycznie za pomocą **współczynnika korelacji**. Ten z definicji przyjmuje wartość pomiędzy 1 a -1³⁷. Wartość dodatnia oznacza pozytywny związek pomiędzy cenami dwóch produktów: gdy cena A rośnie, rośnie też cena B, co wskazuje, że produkty owe należą do tego samego rynku właściwego. Wartość zbliżona do 0 oznacza, że prawie nie istnieje związek pomiędzy cenami produktów. Wartość ujemna oznacza, że ceny reagują w przeciwnych kierunkach i dlatego odnoszą się do odrębnych rynków.

Współczynnik korelacji obejmuje z definicji korelację pomiędzy cenami dwóch produktów. Jeśli bada się, czy więcej produktów należy do tego samego rynku właściwego, obliczyć należy współczynniki korelacji dla wszelkich par produktów.

³⁷ Współczynnik korelacji oblicza się jako: $COV(A,B) / [SD(A)*SD(B)]$. Przy czym $COV(A,B)$ oznacza współczynnik kowariancji między zmiennymi A i B, a $SD(A)$ i $SD(B)$, odpowiednio, odchylenie standardowe zmiennej A i odchylenie standardowe zmiennej B.

Produkty, dla cen których korelacja jest wysoka (a więc blisko 1), mogą należeć do tego samego rynku. Zgodnie z praktyką Komisji Europejskiej, współczynnik korelacji wyższy niż 0,8 kwalifikowany jest jako wysoki, natomiast współczynnik korelacji poniżej 0,65 jako niski³⁸.

Zaletą analizy korelacji cen jest to, że wymogi w stosunku do danych są stosunkowo niewielkie. Bazuje się tu na danych cenowych. Z wysokiej korelacji cenowej pomiędzy produktami A i B nie wynika jednak koniecznie, że produkt B stanowi dla wystarczająco dużej części użytkowników substytut produktu A tak, że mogłoby to hipotetycznego monopolistycznego oferenta powstrzymać od podwyżki o 5-10%, i że w związku z tym produkty należą do tego samego rynku. Innymi słowy, wysoka korelacja cenowa jest niezbędnym, ale niewystarczającym warunkiem uznania dwóch produktów za należące do tego samego rynku, stanowi jednak istotną wskazówkę.

Innym problemem przy analizach korelacji cenowej jest występowanie korelacji pozornych, nieuzasadnionych związkiem konkurencyjnym pomiędzy produktami. Te pozorne korelacje mogą często mieć następujące przyczyny. Kiedy dwa produkty mają **wspólny składnik kosztowy**, to również ich ceny są ściśle skorelowane, pomimo iż produkty te nie są substytutami³⁹. Gdy ceny dwóch produktów są w dużym stopniu zależne od **inflacji** lub **zjawisk sezonowych**, to ich ceny są dodatnio skorelowane, ponieważ rosną wraz z inflacją bądź zmianami sezonowymi, choć jednocześnie nie występuje między nimi faktyczna silna korelacja⁴⁰. Ponadto przedsiębiorstwa z opóźnieniem dostosowują swoje ceny do zmieniających się warunków, tak że dochodzi do **efektów dopasowania**. Jeżeli rozpatrywane interwały czasowe są tak krótkie, że reakcja dopasowania konkurenta zachodzi nie w tym samym, lecz dopiero w następnym okresie, to stwierdzana statystycznie korelacja nie odpowiada korelacji rzeczywistej. Na korelacje pozorne wpływ mogą mieć również

³⁸ Patrz np. Comp (M.2187 CVS Lenzing.

³⁹ Przykładowo, rozwój cen ropy ma wymierny wpływ na rozwój cen benzyny w różnych krajach. Tym niemniej, nie należy wychodzić z założenia, że stacje benzynowe w różnych krajach przynależą do tego samego geograficznego rynku.

⁴⁰ Przykładowo, olej opałowy i warzywa w zimie drożeją, a więc są dodatnio skorelowane, mimo że nie przypisane do tego samego rynku. Pomocniczo zaleca się rozpatrywać rzeczywisty (deflacyjny) rozwój cen. Podobne problemy pojawiają się też z innymi, przeciętnie niestalymi wielkościami (np. PKB).

kursy wymiany walut, kiedy ceny w różnych krajach zależą od wspólnego kursu wymiany.

7 Pozycja dominująca

7.1 *Pozycja dominująca jednego przedsiębiorstwa / indywidualna pozycja dominująca*

Bundeskartellamt: Auslegungsgrundsätze zur Prüfung von Marktbeherrschung in der deutschen Fusionskontrolle, dostępne w Internecie pod adresem: http://www.bundeskartellamt.de/wDeutsch/merkblaetter/Fusionskontrolle/MerkblFusion.shtml
Komisja Europejska: Guidelines on the assessment of horizontal mergers under the Council Regulation on the control of concentrations between Undertakings (2004/C 31/03). dostępne w Internecie pod adresem: http://ec.europa.eu/eurlex/lex/LexUriServ/LexUriServ.do?uri=OJ:c:2004:031:0005:0018:en:PDF .
Department of Justice; Federal Trade Commission: Horizontal Merger Guidelines. dostępne w Internecie pod adresem: http://www.ftc.gov/bc/docs/horizmer.htm .
Kowalski, A. (1997): Die Marktprozessanalyse der Harcard School und neue Systemtheorie, Hamburg, pp. 139 i dalsze.
Vickers, J.; Hay, D. (wydanie) (1987): The Economics of Market Dominance, Oxford.

Pozycję dominującą na rynku można zdefiniować przy pomocy trzech elementów:

1. **Siły ekonomicznej** na danym rynku,
2. Możliwości **zapobiegania** efektywnej **konkurencji**,
3. Możliwości (do pewnego stopnia) **niezależnego postępowania**.

W celu stwierdzenia siły rynkowej można zastosować szereg wskaźników, które zostały przedstawione w poniższych podpunktach. Pozycja rynkowa przedsiębiorstwa nie może być przy tym określona na podstawie zbadania jednego wskaźnika jak np. udział na rynku, lecz na podstawie **całościowego potraktowania wszystkich istotnych okoliczności**. W indywidualnych przypadkach mogłoby to doprowadzić do tego, że pomimo wysokiego udziału w rynku, nie można by zakładać istnienia pozycji dominującej, ponieważ inne czynniki strukturalne, jak np. efektywna

potencjalna konkurencja zagraniczna lub silna przeciwna strona rynku, według wszelkich oczekiwań, nie pozostawia żadnych sfer działania bez kontroli ze strony pozostałych graczy rynkowych.

7.1.1 Udział w rynku

Na wielu rynkach **udział w rynku** stanowi ważną zmienną, pozwalającą na wyciągnięcie wstępnych wniosków odnośnie potencjału danego przedsiębiorstwa i jego możliwości działania w przyszłości. Ocena udziału w rynku ma szczególne znaczenie informacyjne, jeżeli obok **absolutnej wysokości udziału** w rynku ustalone i ocenione zostaną także różnica pomiędzy udziałem w rynku lidera i jego największego konkurenta oraz **rozkład** udziałów w rynku, jak również zmiany powyższych wielkości w trakcie kilku okresów.

Ocena różnicy pomiędzy udziałem rynkowym lidera i jego największego konkurenta oraz rozkładu udziałów w rynku umożliwia wyciągnięcie wniosków dotyczących zdolności konkurentów do zaoferowania przeciwnej stronie rynku własnych produktów, gdyby podmiot dominujący wykorzystywał swą pozycję w sposób y ograniczający konkurencję. Im większa jest różnica pomiędzy udziałem rynkowym lidera i drugiego co do wielkości podmiotu, im bardziej rozdrobnione są udziały w rynku pozostałych konkurentów i im stabilniejsza w miarę upływu czasu jest przewaga lidera rynkowego, tym bardziej prawdopodobna jest, iż przedsiębiorstwo to posiada siłę rynkową umożliwiającą ograniczanie konkurencji. Odnosi się to w jeszcze większym stopniu do rynków, na których działają przede wszystkim przedsiębiorstwa średniej wielkości.

Przy ocenie koncentracji rynku, użyteczne są tak zwane miary koncentracji. Są to ważne wskaźniki stosowane przez organy ochrony konkurencji do rozstrzygnięcia kwestii, czy istnieje dominacja na rynku i czy należy ingerować, aby zapobiec dalszym szkodom dla konkurencji.

Dwa czynniki mają wpływ na koncentrację: po pierwsze liczba przedsiębiorstw w danej branży, a po drugie rozłożenie produkcji pomiędzy te przedsiębiorstwa. Wymiar koncentracji powinien uwzględniać oba czynniki. Poniżej zostaną zaprezentowane dwie istotne dla pracy organów ochrony konkurencji miary koncentracji: udział w rynku i współczynnik Hirschmana-Herfindahla.

Poniżej n oznacza liczbę przedsiębiorstw pewnej branży a Y zagregowaną produkcję ew. zagregowany obrót. Produkcja wzgl. obrót przedsiębiorstwa i oznaczono jako y_i , gdzie i może przyjąć wartości od 1 do n (całkowita liczba przedsiębiorstw na jednym rynku). Przyjmuje się więc

$$Y = \sum_{i=1}^n y_i.$$

Udział w rynku s_i to procentowy udział produkcji/obrotów⁴¹ jednego przedsiębiorstwa w całkowitej produkcji wzgl. całkowitych obrotach tej branży:

$$s_i = \frac{y_i}{Y}$$

Udziały w rynku wszystkich przedsiębiorstw równe są 100%, obowiązują:

$$\sum_{i=1}^n s_i = \frac{\sum_{i=1}^n y_i}{Y} = 1.$$

Z dotychczasowej praktyki decyzyjnej Komisji Europejskiej można wyciągnąć następujące wnioski: podczas gdy przy udziale w rynku poniżej 25% wykluczone jest występowanie indywidualnej pozycji dominującej na rynku, to przy udziale w rynku w przedziale 25-39% jest ona możliwa tylko w rzadkich sytuacjach. Przy udziałach w przedziale 40-49% stwierdzenie dominacji na rynku zależy od siły aktualnej i potencjalnej konkurencji. Udział w rynku powyżej 50% jest znaczącą przesłanką wskazującą na występowanie pozycji dominującej.⁴²

⁴¹ Orzecznictwo pojmuje udział w rynku – szczególnie w przypadku produktów heterogenicznych – z reguły jako wyrażony poprzez obrót wartościowy udział produktu w rynku a nie jako udział ilościowy. Inne metody wyliczenia mogą być w niektórych okolicznościach bardziej miarodajne lub wychwycić trudności w ustaleniu wartości obrotów. Dotyczyło to w przeszłości np. rynków branży budowlanej wzgl. sektora zbrojeniowego (wartość zleceń) lub transportu lotniczego (liczba przewożonych pasażerów). Pozostałe trudności w wyznaczaniu rynku właściwego (np. rynku geograficznego stanowiącego podstawę oceny wielkości rynku) można usunąć stosując progi tolerancji przy obliczaniu wysokości udziałów w rynku.

⁴² ETS przyjął w swej decyzji T-210/01 (GE przeciwko KE) z dnia 14.12.2005 założenie, że przy udziałach w rynku > 50% istnieje domniemanie posiadania pozycji dominującej (por. również orzeczenie ETS w sprawie C-62/86 (AKZO przeciwko KE) z dnia 3.07.1991, p. 60) .

Zmiany udziałów w rynku przez kilka okresów mogą również stanowić cenną wskazówkę odnośnie istnienia lub nieistnienia dominującej pozycji rynkowej. Konkurencja jest dynamicznym procesem wzajemnego prześcigania się przedsiębiorstw. Na rynku, na którym panuje konkurencja, dochodzi najczęściej z **biegiem czasu do wahań udziałów w rynku**. Tak jak utrzymujący się długoterminowo wysoki udział w rynku wskazuje na możliwość działania niezależnie od konkurentów, tak wahania udziałów w rynku – w połączeniu ze zmianami lidera na rynku – lub utrzymujący się silny spadek udziałów w rynku mogą przemawiać przeciwko temu.

W przypadkach, w których produkty, np. **dobrze inwestycyjne o długiej żywotności**, są z reguły przedmiotem popytu w odległych w czasie momentach i w ramach umów długoterminowych, jedynie ocena udziałów w rynku obejmująca długie okresy może przynieść jakiegokolwiek wnioski dotyczące walki konkurencyjnej na tych rynkach.

Przy badaniu kształtowania się udziałów rynkowych w czasie należy brać pod uwagę potencjalne **przyczyny zmian udziałów w rynku**. Często można z tego wyciągnąć ważne wnioski co do podważalności pozycji rynkowej badanego podmiotu. I tak spadek udziałów w rynku przy silnej konkurencji cenowej świadczy o konkurencyjności rynku. To samo dotyczy oczekiwania znaczącego spadku udziałów w rynku na skutek reakcji odbiorców polegającej na zmianie dotychczasowych dostawców. Jeżeli dominant na rynku mógł w miarę upływu czasu utrzymać na tym samym poziomie lub wręcz podnieść swój udział w rynku, jest to mocny argument na niepodważalność jego pozycji rynkowej.⁴³

Koncentracja rynku mierzona jest w typowych sytuacjach za pomocą **wskaźnika Hirschmana-Herfindahla** (HHI). Jest on funkcją udziałów przedsiębiorstwa w rynku, z tego powodu wartości wskaźnika HHI uzależnione są również od różnicy pomiędzy udziałami w rynku poszczególnych uczestników; ponadto mamy tu do czynienia z absolutną sumaryczną miarą koncentracji, ponieważ brane pod uwagę są wszystkie przedsiębiorstwa w danej branży.

⁴³ Bundeskartellamt, decyzja z dnia 20.09.1999 "Henkel/Luhns", Rz 34 f. W niniejszym przypadku Henkel mógł nawet w niesprzyjających okolicznościach konkurencji – zmniejszenie wielkości rynku, wzrost udziałów w rynku producentów marek handlowych kosztem producentów marek najwyższej jakości – zwiększyć swoje udziały na rynku.

Formalnie HHI definiuje się jako:

$$H = \sum_{i=1}^n (100s_i)^2,$$

tj. sumę kwadratów udziałów w rynku wszystkich przedsiębiorstw (w procentach, stąd pomnożone przez 100) w danej branży. Jeżeli występuje tylko jedno przedsiębiorstwo, wskaźnik HHI wyniesie 10.000, jeżeli na rynku występują dwa przedsiębiorstwa posiadające identyczny w nim udział, to HHI wyniesie 5.000, przy czterech przedsiębiorstwach z równym udziałem w rynku - 2.500. W przypadku, gdy dwa z czterech przedsiębiorstw mają udziały po 40% każde, a pozostałe dwa po 10%, to wskaźnik wyniesie

$$HHI = (100 s_1^2) + (100 s_2^2) + (100 s_3^2) + 100 s_4^2 \quad <==>$$

$$HHI = (100*0,4)^2 + (100*0,4)^2 + (100*0,1)^2 + (100*0,1)^2 \quad <==>$$

$$HHI = 40^2 + 40^2 + 10^2 + 10^2 \quad <==>$$

$$HHI = 1600 + 1600 + 100 + 100 \quad <==>$$

$$HHI = 3400.$$

Na tym przykładzie wyraźnie widać, w jaki sposób różnice pomiędzy udziałami w rynku uwzględniane są przez wskaźnik HHI. Przy w pełni równomiernym rozłożeniu udziałów w rynku na cztery przedsiębiorstwa wskaźnik HHI wyniesie 2.500, w przypadku powyższym o asymetrycznym rozłożeniu udziałów w rynku wskaźnik HHI wzrośnie do 3.400.

Przy HHI wynoszącym ok. 1.000 dominacja na rynku jest w praktyce niemalże wykluczona. Od wartości ok. 1.800 koncentracja zaczyna budzić obawy, z tego powodu np. FTC traktuje w tych przypadkach (wartości HHI dla danego rynku na poziomie 1800 lub wyższym) będący konsekwencją połączenia przedsiębiorstw przyrost (tzw. delta) HHI o 50 punktów jako krytyczny.⁴⁴ Komisja Europejska uznaje wysokość HHI i jej zmiany za podstawowe wskaźniki dominacji rynkowej. Przy czym

⁴⁴ por. Section 1.51 Horizontal Merger Guidelines, w internecie online: <http://www.ftc.gov/bc/docs/horizmer.htm>.

powstanie lub wzmocnienie pozycji dominującej na rynku przy HHI wynoszącym od 1000 do 2000 oraz wynikających z łączenia się przedsiębiorstw przyrostach HHI poniżej 250 uważa się za nieprawdopodobne. To samo dotyczy sytuacji, gdy HHI przekracza 2000, a przyrosty są mniejsze niż 150.⁴⁵

7.1.2 Siła finansowa, siła zasobów

Odpowiednio wykorzystana siła finansowa otwiera przed przedsiębiorstwem dodatkowe możliwości strategiczne. Często może ona zostać zastosowana w celu wykluczenia z rynku niepożądanych konkurentów. W razie konieczności możliwy jest transfer zysków lub wyrównanie strat na różnych rynkach (subwencjonowanie skrośne). Efektem takich działań może być faktyczna rezygnacja przez konkurentów z aktywnego konkurowania (użycia parametrów konkurencji, takich jak cena itp.) oraz powstrzymanie się od wejścia na rynek przez potencjalnych konkurentów.

Na rynkach charakteryzujących się intensywnie prowadzoną działalnością badawczo – rozwojową znaczenia mogą nabrać czynniki takie jak posiadanie wykwalifikowanego personelu lub znaczący potencjał innowacyjny.⁴⁶

Siła finansowa przedsiębiorstwa może być oceniona na podstawie wskaźników finansowych takich jak np. Cash Flow⁴⁷, zysk, rezerwy kapitałowe, lub także dostęp do krajowych i międzynarodowych rynków kapitałowych. W praktyce często stosuje się w tym celu trudny do zmanipulowania wskaźnik całkowitych przychodów przedsiębiorstwa, aby określić i porównywać siłę finansową. Daje on jednak jedynie wstępne wskazówki dotyczące siły finansowej.

⁴⁵ por. numer 19 i dalsze Guidelines on the assessment of horizontal mergers under the Council Regulation on the control of concentrations between Undertakings (2004/C 31/03): <http://ec.europa.eu/eur-lex/lex/LexUriServ/LexUriServ.do?uri=OJ:c:2004:031:0005:0018:en:PDF>.

⁴⁶ Bundeskartellamt, uchwała z dnia 20.09.1999 "Henkel/Luhns", Rz. 41 ff. W przypadku "Corning/BICC" Bundeskartellamt zidentyfikował znaczące moce badawczo-rozwojowe u nabywcy, nie stwierdził jednak niebezpieczeństwa powstania w wyniku transakcji pozycji dominującej na rynku niemieckim z uwagi na istnienie na rynku równie silnego pod względem zasobów konkurenta, uchwała z dnia 02.07.1999, str. 15; por. także Bundeskartellamt, uchwała z dnia 03.03.2000 "Cisco/IBM", Rz. 88 ff.

⁴⁷ Cash-Flow = wynik finansowy brutto + odpisy amortyzacyjne + przyrost (-zmniejszenie) długoterminowych rezerw.

7.1.3 Dostęp do rynków zakupu i sprzedaży

Lepszy w stosunku do konkurentów dostęp do rynków zakupu i zbytu może stworzyć przedsiębiorstwu wiodącą pozycję na rynku. Ma to znaczenie zwłaszcza wtedy, gdy przedsiębiorstwo mające silną (udziałowo) pozycję rynkową z uwagi na swój dominujący dostęp do rynków zakupu i zbytu może utrudniać lub wręcz zamykać dostęp do tych rynków swoim konkurentom (podwyższenie barier dostępu do rynku).

Efekt zamknięcia rynku często występuje w praktyce wtedy, gdy przedsiębiorstwo działa nie tylko na danym rynku, lecz równocześnie na jednym z ważnych rynków poprzedzających lub następujących w łańcuchu produkcji (**integracja pionowa**) i na obu rynkach zajmuje co najmniej silną pozycję.

W niektórych przypadkach różnorodność asortymentu, wielkość produkcji i istniejące struktury sprzedaży wchodzą w miejsce siły finansowej jako jeden z głównych parametrów konkurencyjności. Porównywalne do efektu zamknięcia rynku, chociaż najczęściej mniej sztywne ograniczenie konkurencji, może być spowodowane **oferowaniem szerokiego asortymentu**. Dotyczy to przede wszystkim rynków, na których ogromną rolę odgrywają kontakty z klientem oraz zaufanie użytkowników do uznanej i szerokiej oferty produktów, o ile konkurenci oferenta asortymentu są w stanie zaoferować uboższy asortyment towarów lub usług lub w ogóle nie mają takiej oferty. To samo dotyczy – przy odpowiednim popycie – **oferty kompletnych systemów**, jeżeli konkurenci oferują jedynie komponenty systemu a nie dysponują porównywalną „zdolnością oferowania systemu“.

Lepszy dostęp do rynków zakupu lub rynków zbytu z powodu korzyści konkurencyjnych wynikających z zasobów przedsiębiorstwa może wynikać przykładowo z poniższych czynników:

- Duże poważanie lub szczególne znaczenie rynku / marki
- Własna sieć filii o dużym zagęszczeniu, sprawdzona logistyka sprzedaży, wyróżniająca się obecność terytorialna
- Dostęp do ważnych możliwości transportowych przy ofercie usług dystrybucyjnych (zabezpieczony w umowie spółki)

- Zabezpieczone zakupy półproduktów, ew. (zabezpieczone prawnie) powiązania z producentem na rynku podrzędnym lub nadrzędnym w łańcuchu produkcji.

Prawdopodobieństwo, że przedsiębiorstwo posiadające największy udział w rynku zajmie na tym rynku pozycję dominującą, może zostać zwiększone w efekcie przewagi nad konkurentami wynikającej z posiadanych zasobów.

7.1.4 Powiązania

Powiązania przedsiębiorstwa z innymi, szczególnie konkurentami, odbiorcami lub dostawcami, mogą przyczyniać się do wiodącej pozycji na rynku, ale nie mogą stanowić jedyne jej powodu. O ile istniejące powiązania prawne przedsiębiorstw mogłyby mieć wpływ na ich pozycję rynkową, należy wziąć je również pod uwagę dokonując oceny ogólnej. Nie jest konieczne, by powiązane ze sobą przedsiębiorstwa z punktu widzenia konkurencji występowały jako jeden organizm gospodarczy, ponieważ stosunki ekonomiczne, prawne i handlowe pomiędzy przedsiębiorstwami mogą mieć wpływ na pozycję rynkową nawet wtedy, gdy nie stwierdza się ani wspólnych, ani jednolitych działań rynkowych.

Zasługujące na uwagę powiązania nie muszą mieć **charakteru stosunku prawnego**; mogą one być **natury osobistej, prawnej lub ekonomicznej**. Stosunki kontraktowe – jak np. wzajemne umowy o dopuszczeniu stosowania patentów lub (wyłącznościowe) kontrakty na dostawy – mogą wywoływać efekty dławiące konkurencję. Wszystkie wymienione dotąd rodzaje powiązań z natury mogą wzmocnić pozycję dominującą na rynku, jeżeli przedsiębiorstwo zwiększy w ten sposób swoje możliwości wpływu lub posiadane zasoby. Wynikające z powiązań ryzyka i zagrożenia dla konkurencji zostaną w dużej części poruszone w kontekście innych kryteriów omawianych w tym rozdziale. Tak więc powiązania z przedsiębiorstwami silnymi finansowo mogą zwiększyć uwarunkowane finansowo możliwości działania przedsiębiorstwa, a powiązania z odbiorcami lub dostawcami mogą zabezpieczyć zbyt lub zakupy oraz podnieść bariery dostępu do rynku

Szczególnie duże znaczenie ma to w przypadku **powiązań z aktualnymi lub potencjalnymi konkurentami** oraz z oferentami niedoskonałych substytutów, ponieważ poprzez to ograniczona zostaje zazwyczaj konkurencja pomiędzy konkurentami. Powiązania poprzez wspólne przedsiębiorstwo produkujące

półprodukty wzgl. komponenty działają jednocząco na warunki zbytu konkurentów z powodu ujednolicenia podstawy kosztów. To samo dotyczy w jeszcze większym stopniu współpracy pomiędzy konkurentami na właściwym rynku produktowym.

7.1.5 Bariery wejścia / potencjalna konkurencja

Sprawdzenie barier wejścia na rynek i potencjalnej konkurencji ma duże znaczenie w ocenie pozycji rynkowej. Bariery wejścia dają informację o znaczeniu potencjalnej konkurencji dla stosunków konkurencyjnych na badanym rynku. Dostęp do rynku jest w tym znaczeniu kryterium odnoszącym się nie do przedsiębiorstwa, lecz dotyczącym struktury rynku. Dopóki silne przedsiębiorstwo na rynku nie żąda wygórowanych cen lub nie może zrezygnować z badań i rozwoju, ponieważ w innym wypadku musi się liczyć z wejściem na rynek potencjalnych konkurentów, jest mało prawdopodobne, aby posiadało niepodlegające istotnym ograniczeniom możliwości postępowania.

Wysokie bariery wejścia na rynek mogą być istotną przesłanką wskazującą na pozycję dominującą przedsiębiorstwa, ponieważ zabezpieczają to przedsiębiorstwo przed nowymi wejściami na rynek. Przy badaniu kwestii potencjalnej konkurencji należy zbadać, czy **możliwe i prawdopodobne** byłoby znaczące z punktu widzenia intensywności konkurencji wejście na rynek. Musi ono być do tego wystarczająco **realne**. Często nie ma zachęty do wchodzenia na sąsiadujące rynki przy w pełni wykorzystanych mocach i ugruntowanych związkach z klientami. Ponadto rodzi się pytanie, czy przedsiębiorstwo może wejść na rynek z takimi **ilościami i cenami**, które będą w stanie skutecznie zawęzić niepodlegający istotnym ograniczeniom margines działań rynkowych.

Teoria spornych rynków (contestable markets) dostarcza w ramach ekonomii przemysłu istotnych wniosków. Odpowiednie uwarunkowania dla takich modeli rzadko można jednak znaleźć w rzeczywistości. W szczególności wysokość barier dostępu do rynku, jako czynnik wpływający na realność zagrożenia wejściem na rynek potencjalnych konkurentów, należy zbadać w każdym pojedynczym przypadku na podstawie warunków panujących na rynku. Nadzorując konkurencję należy dążyć do połączenia oceny aktualnej i potencjalnej konkurencji. Ogromną rolę odgrywa przy tym sprawdzenie rzeczywistych „możliwości zaatakowania” badanego rynku – w szczególności odnosi się to do fazy rynku, barier wejścia, potencjalnej konkurencji.

Rynek jest tym bardziej sporny, im bardziej bezproblemowe jest dojście do rynków zbytu wzgl. rynków zaopatrzenia (patrz 7.1.3) lub do najlepszej technologii (7.1.2), im mniejsze są związane z wyjściem z rynku koszty utopione (sunk costs) oraz im wcześniej potencjalny oferent może zakładać, że podmiot o pozycji dominującej na rynku w przypadku nowego wejścia utrzyma swoją cenę odpowiednio długo. Mimo wszystko w rzeczywistości nawet w takich warunkach występują mniej lub bardziej wyraźne bariery wejścia. Ponieważ wejście na nowy rynek często wymaga czasu, wchodzenie odbywa się najczęściej krok po kroku. Im dłuższy jest **czas potrzebny na wejście na rynek**, tym bardziej należy się liczyć z reakcją oferentów obecnych już na rynku. Ponadto nowi przedsiębiorcy chcąc **pozyskać kapitał zewnętrzny** muszą, jak uczy doświadczenie, liczyć się z wyższymi odsetkami niż przedsiębiorstwa działające już na rynku, gdyż stanowią większe ryzyko dla kapitałodawcy.

Wysokie bariery wejścia nie muszą w pełni wykluczać nowych wejść na rynek, aby można było stwierdzić istnienie pozycji dominującej. Bliższe prawdy jest stwierdzenie, że oczekiwane wejście na rynek nie będzie miało **takich rozmiarów**, które byłyby w stanie ograniczyć swobodę postępowania danego przedsiębiorstwa. Do tego wejście na rynek musi nastąpić w **ramach czasowych**, które będą na tyle krótkie, aby badane przedsiębiorstwo nie mogło wykorzystać swojej wiodącej pozycji rynkowej. Nieproporcjonalnie **wysokie koszty wejścia na rynek lub duże ryzyko błędnej decyzji** są przesłankami wysokich barier dostępu do rynku. Należy przyjmować takie założenie szczególnie w przypadku **subaddytywnej struktury kosztów** w połączeniu z wielkimi, **nieodwracalnymi inwestycjami**, które są konieczne celem wejścia na rynek, a w przypadku wyjścia z rynku stanowiłyby koszty utopione.

W zależności od ich pochodzenia można podzielić **bariery dostępu do rynku** na trzy grupy:

1. **Prawne** bariery dostępu do rynku tworzone są w oparciu o państwowe regulacje w formie ustaw, rozporządzeń lub praktyki administracyjnej.
2. **Strukturalne** bariery dostępu do rynku tłumaczy się z reguły technologiczną lub uwarunkowaną popytem charakterystyką rynku, ale mogą one też polegać

na warunkującej sukces rynkowy przedsiębiorstwa sile jego zasobów. Najczęściej nie są one tworzone z zamiarem ochrony rynku przed wejściem.

3. **Strategiczne** bariery dostępu do rynku z kolei tworzone są przez aktualnych uczestników rynku świadomie, aby odstraszyć potencjalnych oferentów od wchodzenia na rynek.⁴⁸

Nowych wejść na rynek należy oczekiwać tym bardziej, im wyżej można ocenić perspektywę przyszłych zysków. Nowe i rozrastające się rynki lub rynki z nadwyżką popytu wykazują najczęściej mniejsze bariery wejścia niż rynki dojrzałe z nadwyżką mocy produkcyjnych. Jeżeli czas pozostający na amortyzację nowych inwestycji jest krótki, ponieważ rynek ma charakter zanikający, wejście na rynek jest mało prawdopodobne.

7.1.6 Konkurencja poprzez substytucję brzegową (marginalną)

Swoboda postępowania przedsiębiorstw dominujących na rynku może być w ograniczonym zakresie zmniejszona przez przedsiębiorstwa, które oferują towary i usługi niebędące wprawdzie takiej samej wartości jak towary rynkowe, ale mogące je zastępować w ograniczonym zakresie lub w pewnych okolicznościach. Istnienie spokrewnionych rynków, których produkty zaspakajają podobne potrzeby, może przyczynić się do ograniczenia dominacji na rynku, ponieważ istnieją możliwości ominięcia dominanta. Przykładem może być tutaj branża transportowa: w zależności od charakterystyki towarów przeznaczonych do transportu (wrażliwość na czas, ciężar itp.) i trasy przewozu nabywcy mają do wyboru różne środki transportu jak autobus, samochód, samolot, okręt lub kolej (tzw. konkurencja intermodalna). Tego typu **konkurencja substytucyjna, bądź „resztowa”** na rynku charakteryzującym się obecnością przedsiębiorstwa o pozycji dominującej nie zapewnia jeszcze sama z siebie wystarczającej kontroli swobody działania dominanta rynku.

Niedoskonała substytucyjność może wynikać z różnych powodów. Przykładowo zmiana na (niedoskonałe) substytuty może być **możliwa jedynie długoterminowo długim okresie**, ponieważ najpierw należy dokonać poważnych inwestycji, aby

⁴⁸ W szczególności szkoła harwardzka podkreśla w analizie barier dostępu do rynku i potencjalnej konkurencji istnienie strategicznych barier wejścia na rynek, które mogą polegać na przykład na inwestycjach obejmujących koszty utopione (sunk costs).

można je w ogóle zastosować. Dotyczy to przykładowo przestawienia się z oleju opałowego na gaz jako nośnik energii cieplnej. Konkurencja odbywa się tu głównie w fazie otwierania (fikcyjnego) rynku energii cieplnej, ponieważ w tym czasie podejmowane są decyzje o podstawowych inwestycjach związanych z ułożeniem instalacji.

Wybór niedoskonałego substytutu może wynikać ze **zmiany gustów lub przyzwyczajeń nabywców**, na przykład przy zmianie z tygodnika politycznego na ogólnokrajową gazetę codzienną. Ponadto zmiana taka może być utrudniona z uwagi na różnicę **cen** lub sposób użycia. Równowartość rynkowa nie występuje także wtedy, gdy pomimo nakładania się zakresów zastosowania użycie substytutu wymaga wyższych **nakładów czasowych i pieniężnych**. Nakładanie się jedynie częściowo zakresów zastosowania nie wystarcza z reguły do założenia równowartości rynkowej produktów i usług.

Konkurencja substytucyjna może pochodzić od wielu towarów i usług, które przyporządkowane są różnym rynkom produktowym. Intensywność konkurencji może być przy tym bardzo różna, w zależności od tego, jak dobrze z punktu widzenia nabywcy jeden towar może zostać zastąpiony przez inny. W przypadku **towarów lub usług heterogenicznych** intensywność konkurencji substytucyjnej może być także bardzo różna dla poszczególnych oferentów na rynku pierwotnym.

7.1.7 Przeciwważna siła rynkowa

Wysoka koncentracja po stronie popytu nie jest sama w sobie wystarczającą przesłanką aby podważyć dominującą pozycję na rynku jednego z oferentów, ponieważ jakakolwiek **siła rynkowa po stronie popytu** odnosi się do wszystkich **oferentów w równej mierze**. Wymogiem dla zaistnienia przeciwważnej siły rynkowej jest raczej okoliczność, strategicznego rozdzielenia zleceń przez posiadającego siłę rynkową nabywcę, w celu uniknięcia uzależnienia od (dominującego) dostawcy. Takie strategiczne postępowanie jest wiarygodne, gdy rynek jest szczególnie ważny dla kupujących na nim przedsiębiorstw, np. ponieważ koszty półproduktu decydują w znacznej mierze o cenie sprzedaży wytwarzanego produktu ostatecznego. O sile rynkowej po stronie popytu można mówić także wtedy, gdy zakupy nabywcy dotyczą **większej części produkcji na rynkach ofertowych**. Dla silnego na rynku oferenta, który większą część swojej produkcji dostarcza jednemu nabywcy, zagrożenie

zmiany dostawcy – o ile konkurenci dysponują wystarczająco dużymi mocami produkcyjnymi – może mieć skutki silnie dyscyplinujące.

Przy ocenie przeciwważnej siły rynkowej rodzi się pytanie, czy funkcje efektywnej konkurencji – takie jak skuteczne ograniczanie kosztów i cen oraz zachęta do postępu technicznego – w indywidualnym przypadku mogą zostać w sposób równoważny zastąpione siłą popytu oraz strategicznymi działaniami w zakresie zaopatrzenia. Odpowiedź na tak sformułowane pytanie mogłaby być twierdząca, jeśli chodzi o przedsiębiorstwa o dużej sile po stronie popytu w stosunku do zoligopolizowanej podażowej strony rynku (czyli siła po stronie popytowej może w niektórych przypadkach neutralizować negatywne konsekwencje oligopolu po stronie podaży). Trudniej jest odpowiednio udowodnić relacje pomiędzy przedsiębiorstwami o dużej sile popytowej a pojedynczym oferentem o dużej sile rynkowej.

Utrzymanie istotnych funkcji konkurencji poprzez strategiczne postępowanie w zakresie zaopatrzenia wymaga, aby nabywcy dysponowali szerokim zakresem **informacji** w odniesieniu do danych produktów. To z kolei uzależnione jest m.in. od rodzaju towarów, będących przedmiotem wymiany.⁴⁹

Specyfika rynku może ograniczać znaczenie siły rynkowej po stronie popytu. I tak strategiczne postępowanie nabywcy wykluczone jest wtedy, gdy przeciwna strona rynku nie podejmuje sama decyzji o zakupie niniejszych produktów i własna polityka zakupów w tych okolicznościach możliwa jest tylko w ograniczonym stopniu. **Siła popytowa handlu** nie jest wystarczającą przeciwwagą wobec dominującej pozycji po stronie producenta, jeżeli dany produkt jest wiodący na rynku, z uwagi na kosztowną reklamę oraz silne przyzwyczajenie konsumenta „sprzedawany jest priorytetowo” i w związku z tym musi znaleźć się na półkach każdego sprzedawcy detalicznego. Groźba, że handel zrezygnuje z wiodącej marki jest szczególnie mało wiarygodna, kiedy producent dysponuje kilkoma lub całym asortymentem wiodących marek najwyższej klasy.

Siła rynkowa po stronie popytu może wykluczyć dominację rynkową po stronie podaży, gdy groźby wszystkich liczących się nabywców są w porównywalnym

⁴⁹ por. w tym temacie zapisów dotyczących zawodności rynku z powodu braku informacji .

stopniu wiarygodne/mocne. Siła rynkowa po stronie popytu nie prowadzi ponadto do skutecznego ograniczenia dominacji rynkowej, kiedy - przykładowo z uwagi na wysoki stopień znajomości na rynku marek dominanta – działa jednostronnie na niekorzyść konkurentów, a nie na niekorzyść podmiotu dominującego.

7.1.8 Faza rynku

Faza rynku mówi o tym, w jakim stanie rozwoju znajduje się dany rynek i dalej o trwałości warunków konkurencji na nim. Nie jest to w tym znaczeniu czynnik samodzielny, lecz pozwala on na wyciągnięcie ważnych wniosków przy ocenie udziałów w rynku i ich zmian w ramach „cyklu życia” rynku.

7.2 Zbiorowa pozycja dominująca

W poprzednim punkcie przedstawiono kryteria stanowiące o indywidualnej pozycji dominującej. Pojawiające się unilateralne (nieskoordynowane) efekty są tym większe, im większa jest siła rynkowa badanego przedsiębiorstwa. Siłę rynkową mogą wykorzystywać nie tylko pojedyncze przedsiębiorstwa, ale także w warunkach oligopolu grupa złożona z wielu większych przedsiębiorstw wspólnie. Efekt unilateralny, który wynika z dominacji na rynku każdego z tych przedsiębiorstw, jest wprawdzie niewielki, ale oprócz efektu unilateralnego może dochodzić na rynku oligopolistycznym także do **skoordynowanego postępowania oligopolistów**. Jeżeli występuje tego typu skoordynowane postępowanie, mowa jest o kolektywnej (również oligopolistycznej) pozycji dominującej.

Dygresja: Postępowanie skoordynowane w ramach teorii gier

Skutki postępowania skoordynowanego można przedstawić na przykładzie konkurencji cenowej w punkcie dotyczącym teorii gier. W przedstawionym tam modelu dla oferentów byłoby korzystne, gdyby porozumieli się wspólnie co do wysokiej ceny i wspólnie osiągnęli wysokie zyski. W warunkach konkurencji jednak okazuje się, że każdy oferent ma indywidualną zachętę do tego, aby oferować cenę niższą od konkurentów. W efekcie konkurencji decydują się krótkoterminowo na taką niekooperatywną strategię, ponieważ muszą liczyć się z tym, że ich konkurenci, każdy z osobna, ustalą niższą cenę i oni sami poniosą w ten sposób stratę.

Przy rozwiązaniu gry wzięto pod uwagę tylko jedną pojedynczą sytuację decyzyjną. W realnej sytuacji rynkowej uczestnicy trafiają jednak wielokrotnie na siebie zawsze w tej samej konstelacji. Jeżeli **gra** odbywa się **powtórnie**, to wynik równowagi może się zmienić i zostać zastąpiony przez wynik kooperatywny, chociaż struktura gry pozostaje taka sama.

Można to przybliżyć pokazując kalkulację decyzyjną jednego z uczestników: Jeżeli jeden z uczestników ustali wysoką cenę, grozi mu niebezpieczeństwo poniesienia straty, ponieważ

jego konkurent może ustalić niższą cenę. Jeżeli jednak można stwierdzić wiarygodnie, że będzie można zastosować sankcje w przypadku ustalenia niższej ceny, to można wyjść z założenia, że konkurent odstąpi od możliwości ustalenia niższej ceny i też ustali wysoką cenę. Taki mechanizm sankcji może polegać przykładowo na tym, że na określony czas ustala się wyjątkowo niską cenę, aby spowodować straty u tego konkurenta. Konkurent musi wtedy wyważyć krótkoterminowy wysoki zysk wynikający z niskiej ceny wobec grożących mu w przyszłości strat. Sankcje muszą być w każdym razie wiarygodne, tj. trzeba mieć zachętę, aby te sankcje rzeczywiście wprowadzić w życie. Wynik w takiej sytuacji odpowiadałby wynikowi kooperacyjnemu, jednak bez rzeczywistego uzgodnienia wzajemnych stanowisk. W literaturze ekonomicznej mowa jest dlatego o **milczącej lub dorozumianej zмовie** (tacit or implicit collusion).

Skoordynowane działanie może dotyczyć różnych parametrów konkurencji. Dla przykładu oligopolisci mogą skoordynować swoje postępowanie w odniesieniu do cen. Podobnie może zaistnieć koordynacja w odniesieniu do wielkości produkcji lub mocy produkcyjnych.

Czy dochodzi jednak do efektów kooperacyjnych, uzależnione jest od wielu różnych czynników. Komisja Europejska określiła trzy niezbędne wymogi na stwierdzenie możliwości występowania efektów koordynacji:

1. Rynek musi być **wystarczająco przejrzysty**. Tylko w takim przypadku przedsiębiorstwa mogą stwierdzić, czy konkurenci będą się trzymać milczącej zмовy.
2. Musi występować **wiarygodny mechanizm sankcji**.
3. Aktualni lub potencjalni konkurenci spoza oligopolu nie są w stanie podważyć uzgodnionego postępowania oligopolistów, tj. nie istnieje **wystarczająca konkurencja zewnętrzna**.

Występują jednak dalsze czynniki, które utrudniają lub sprzyjają koordynacji postępowania. Jeżeli **liczba członków oligopolu** jest relatywnie duża, utrudnia to koordynację, ponieważ członkowie muszą w sposób w pełni zrozumiany zgodzić się co do ceny (lub ilości) i ustalić mechanizm sankcjonujący. Jeżeli oligopolisci mają różną wielkość, czyli są bardzo asymetryczni, to małemu oligopoliscie będzie niezwykle trudno stosować wiarygodny mechanizm sankcji. Także koordynacja w przypadku produktów heterogenicznych jest trudniejsza niż przy produktach homogenicznych. Na rynkach, na których bardzo rzadko dochodzi do transakcji, koordynacja jest trudniejsza niż na rynkach, gdzie oferenci i nabywcy spotykają się często. Koordynacja jest również utrudniona przy dużej elastyczności **popytu**. W

takim przypadku przedsiębiorstwa, ograniczając podaż, mogą przeforsować jedynie niewielką podwyżkę cen, a w ten sposób zyski ze skoordynowanego postępowania będą niewielkie. Także silnie wahający się popyt utrudnia koordynację. Przy wyraźnym obniżeniu popytu przedsiębiorstwo nie jest w stanie bezpośrednio stwierdzić, czy spadek wywołany jest przez wyłamanie się jednej z firm z kolektywu dominującego na rynku. Jeżeli wtedy zostanie zastosowany mechanizm sankcji, to koordynacja rozpadnie się, pomimo iż nie nastąpiło żadne odstępstwo. Jeżeli krótkoterminowo nastąpi silny wzrost popytu, to wzrośnie zysk wynikający z wyłamania się, a w ten sposób mechanizm sankcji straci skuteczność.

Jeżeli z kolei relacje rynkowe są przejrzyste, oligopoliści mają podobną wielkość, popyt jest raczej nieelastyczny i stały, a konkurenci zewnętrzni nie posiadają odpowiednich mocy produkcyjnych, aby wywierać nacisk na oligopolistów, to wzrasta prawdopodobieństwo koordynacji postępowania.⁵⁰ Szerzej kwestie te zostaną przedstawione w dotyczącym tzw. karteli hardcore rozdziale 10.1.2.1.

8 Kontrola koncentracji

8.1 Połączenia horyzontalne

Fuzje horyzontalne oznaczają łączenie się konkurentów, którzy działają na tym samym rynku. Tego typu fuzja może budzić wątpliwości z dwóch różnych powodów: może prowadzić do 1. efektów unilateralnych (opanowanie rynku przez przedsiębiorstwo powstałe w wyniku fuzji) lub 2. efektów koordynacji (powstanie zbiorowej pozycji dominującej).

8.1.1 Efekty unilateralne

Kiedy na rynku dochodzi do fuzji, wtedy zmniejsza się liczba oferentów. Jak pokazano już w ramach teorii oligopolu (rozdz. 4.5) zmniejszająca się liczba

⁵⁰ Więcej szczegółów na temat koordynacji zachowań rynkowych oligopolistów można znaleźć w wytycznych Komisji Europejskiej dotyczących oceny porozumień poziomych. Patrz również decyzja Trybunału Pierwszej Instancji w przypadku *Airtours/First Choice* (T-342/99), w której uchylono decyzję Komisji w sprawie *Comp/M.1524 Airtours/First Choice*, zakazującą połączenia.

oferentów – o ile na skutek fuzji nie następuje wzrost efektywności – najczęściej idzie w parze ze zmniejszeniem wielkości sprzedaży i wzrostem ceny. Fuzja prowadzi w takim przypadku do **spadku dobrobytu** (wyrażonego przez większą stratę społeczną (deadweight loss)).⁵¹ Jeżeli na skutek fuzji przedsiębiorstwo – staje się dominującym na rynku, może podkreślić swoją dominację na rynku przez wzrost ceny, pogorszenie jakości produktów, zmniejszenie nakładów na badania i rozwój itp. Wysokość związanego z tym spadku dobrobytu uzależniona jest od dominacji rynkowej przedsiębiorstwa powstałego na skutek fuzji. Im większa dominacja na rynku, tym mniejsze są rynkowe ograniczenia swobody działania dla przedsiębiorstwa powstałego na skutek fuzji i tym drastyczniej może ono podnosić ceny (lub obniżać jakość itp.) po przeprowadzeniu fuzji. Ponadto skuteczność praktyk monopolistycznych związanych z nadużywaniem pozycji dominującej jest dla przedsiębiorstwa tym większa, im wyższa jest jej siła rynkowa. Wskaźnikami siły rynkowej są w szczególności czynniki⁵² wymienione w części 7.1

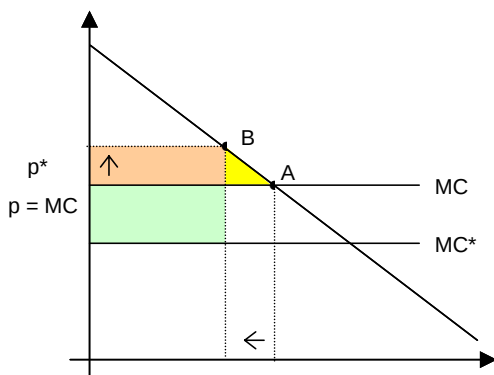
Nieco mniej jednoznaczna ocena efektów dobrobytu ma miejsce wtedy, gdy fuzja w przypadku badanego przedsiębiorstwa prowadzi do **wzrostu efektywności**. Przyczyny wzrostu efektywności mogą być bardzo różne. Mogą to być np. **przychody z korzyści skali** (economies of scale) lub **korzyści zakresu** (economies of scope)⁵³ odnoszące się do produkcji. Przedsiębiorstwo powstałe w wyniku fuzji może również osiągnąć efekty **synergii** w dziedzinie badań i rozwoju, reklamy i dystrybucji, lub poprzez redukcję kosztów administracyjnych. Podczas gdy spowodowany fuzją przyrost siły rynkowej zmniejsza dobrobyt, to przyrost efektywności **dobrobyt podnosi**. W takiej sytuacji możliwe jest, że pozytywny efekt wzrostu efektywności przeważa negatywny efekt dominacji rynkowej. Poniższy wykres przedstawia to zjawisko.

⁵¹ Osiągnięcie pozycji dominującej nie jest jednak koniecznym wymogiem do tego, aby fuzji towarzyszył negatywny efekt spadku dobrobytu. Jak pokazano w ramach teorii oligopolu, spadek liczby oferentów prowadzi do podniesienia ceny w konkurencji Cournota lub konkurencji Bertranda przy dobrach heterogenicznych, a co za tym idzie, do spadku dobrobytu.

⁵² W związku z tym warto wspomnieć, że – o ile nie dochodzi do wzrostu efektywności - konkurenci korzystają na władzy rynkowej dominującego przedsiębiorstwa w ten sposób, że sami też mogą podnieść ceny lub wielkość produkcji wprowadzaną na rynek.

⁵³ Porównaj w tym temacie szczegóły osobnego rozdziału.

Wykres 24: Wzrost wydajności a zwiększenie siły rynkowej



Przyjęto założenie, że przed przeprowadzeniem fuzji na rynku działało dwóch konkurentów, którzy poprzez fuzję stworzyli monopol. Przed przeprowadzeniem fuzji konkurowali oni ze sobą cenami osiągając w wyniku cenę na poziomie kosztów krańcowych ($p = MC$) (równowaga rynkowa w punkcie A). Po dokonaniu fuzji przedsiębiorstwo straci konkurencję i ustali taką cenę, która znajduje się powyżej kosztów krańcowych. Jednocześnie wraz z fuzją na skutek wzrostu efektywności spadły koszty krańcowe (istotne dla ustalenia ceny) z MC do MC^* , więc po fuzji ustalona została cena (monopolistyczna) „ p^* ” (równowaga rynkowa w punkcie B). Czy teraz cena „ p^* ” będzie wyższa lub niższa od poprzednich kosztów krańcowych, uzależnione jest od skali obniżenia kosztów. Im wyższa obniżka kosztów, tym niższa cena po przeprowadzeniu fuzji.⁵⁴ Chociaż przedstawione przemyslenia bazują na założeniu, że w wyniku połączenia konkurentów powstaje monopol, taki trade off powstanie także w innych układach rynkowych, choćby przy konkurencyjnym oligopolu. Im mniejsza dominacja rynkowa przedsiębiorstwa powstałego na skutek fuzji, tym niższa będzie cena ustalona po obniżeniu kosztów. Fuzja może w związku

⁵⁴ Aby dobrobyt wzrósł, nie jest konieczne, aby cena po fuzji była niższa od tej sprzed połączenia. Wzrost dobrobytu występuje wówczas, gdy dochód producenta osiągnięty dzięki poprawie efektywności (ciemny obszar) jest większy od straty dobrobytu wynikającej z podniesienia ceny (jasny obszar). To, że na skutek fuzji powstaje efekt zmiany podziału dochodu konsumenta i producenta, nie ma znaczenia z punktu widzenia oceny dobrobytu. Jednak zagadnienie, czy taką zmianę podziału należy w ramach oceny skutków połączenia uznać za pozytywną czy negatywną, jest kwestią normatywną. Efekt zmiany podziału nie jest istotny, gdy spadek kosztów jest na tyle duży, że fuzja prowadzi do obniżenia ceny.

z tym już przy małym wzroście efektywności wpłynąć na obniżenie ceny i skutkować wzrostem dobrobytu.

Uwzględnienie przyrostu efektywności podczas kontroli fuzji przysparza pewnych trudności. Po pierwsze nie każdy wzrost efektywności przenosi się na ustalenie ceny. Podczas gdy przychody wynikające z korzyści skali wpływają na zmienne koszty produkcji, to pozostałe korzyści wpływają tylko na **koszty stałe**. Te ostatnie nie mają jednak wpływu na ustalenie ceny przez przedsiębiorstwo powstałe w wyniku fuzji, więc nie dojdzie do obniżenia ceny i lepszego zaopatrzenia rynku w wyniku wzrostu efektywności.⁵⁵ Do tego dochodzi jeszcze fakt, że efektywność techniczna w produkcji *ex ante* da się łatwiej wyliczyć niż np. redukcja (stałych) kosztów administracyjnych. Po drugie **korzyści efektywności** powinny **być brane pod uwagę** tylko wtedy, gdy mogą być one zrealizowane **wyłącznie poprzez fuzję** i nie byłyby osiągalne bez fuzji. Osiąganie efektów skali jest możliwe do zrealizowania także bez fuzji przykładowo poprzez wewnętrzny wzrost, a synergia poprzez redukcję kosztów administracyjnych. Po trzecie istnieje jeszcze jeden problem przy **ocenie wysokości wzrostu efektywności** występuje asymetria informacji pomiędzy organem antymonopolowym a badanym przedsiębiorstwem. Jeśli taka ocena jest istotna przy ocenie przypadku fuzji przez właściwy organ, występuje zachęta dla badanych przedsiębiorstw, aby zbyt wysoko oceniać oczekiwany przyrost efektywności. Dalszym problemem może być w końcu także kwantyfikacja przyrostu efektywności.

8.1.2 Efekty skoordynowane

Fuzja konkurentów może prowadzić nie tylko do opanowania rynku przez jeden podmiot, ale także do kolektywnego opanowania rynku przez jedno lub więcej spośród przedsiębiorstw uczestniczących w koncentracji. Jak już wspomniano, efekty unilateralne fuzji są w takim przypadku mniejsze niż w przypadku opanowania rynku przez jeden podmiot. Jednak może dojść do koordynacji zachowań pomiędzy badanymi przedsiębiorstwami. Jeżeli fuzja stwarza na rynku warunki, które ułatwiają koordynację zachowań pomiędzy konkurentami, należy wychodzić z założenia, że

⁵⁵ Mimo wszystko redukcja kosztów stałych może spowodować pozytywny efekt dla dobrobytu. Wyraża się on jednak tylko w formie wyższego dochodu producenta. Czy należy brać to pod uwagę oceniając konkurencyjność, pozostaje decyzją normatywną.

fuzja ma **skutki ograniczające konkurencję**. Jednym z powodów jest fakt, że poprzez fuzję zmniejsza się liczba przedsiębiorstw uczestniczących w rynku a **koordynacja postępowania** staje się **łatwiejsza**. Oprócz tego fuzja może jednak wzmocnić **symetrię** pomiędzy przedsiębiorstwami i w ten sposób zwiększyć skuteczność ewentualnego mechanizmu sankcji. To, czy poprzez fuzję rzeczywiście dojdzie do kolektywnego opanowania rynku, uzależnione jest od dalszych sprzyjających czynników, które zostały szczegółowo opisane wyżej, są to choćby przejrzystość rynku, siła konkurencji zewnętrznej lub jednorodność produktów.

Również w ramach zbiorowej pozycji dominującej należy brać pod uwagę **przyrost efektywności** spowodowany fuzją. Obowiązują tutaj takie same podstawy i zastrzeżenia, jak w ramach pojedynczej pozycji dominującej. Ponadto istnieje możliwość, że poprzez wzrost efektywności zostanie zakłócona w pewnym stopniu symetria pomiędzy oferentami. Jeżeli np. dokonuje się fuzja pomiędzy członkiem oligopolu z przedsiębiorstwem, które dysponuje zdecydowanie tańszą technologią produkcji, to może być to uzasadniona zachęta dla tego przedsiębiorstwa, aby wyrwać się z oligopolu. Z drugiej strony wzrost efektywności może wyrównać wyższe koszty członka oligopolu i w ten sposób wzmocnić symetrię pomiędzy oferentami.

8.1.3 Modele symulacyjne fuzji

Modele symulacyjne fuzji stwarzają możliwość arytmetycznej oceny korzyści (tylko) unilateralnych efektów fuzji. Ocena efektów koordynacji nie jest możliwa na podstawie tego modelu. Opierając się na danych można stworzyć prognozy oczekiwanych zmian cen i konsekwencji fuzji dla dobrobytu społecznego. Jako dane bazowe można przyjąć na przykład udziały w rynku przedsiębiorstw łączących się i ich konkurentów, elastyczność cenową i mieszaną popytu, obrót na rynku, koszty krańcowe i, o ile występują, także osiągalne oszczędności kosztów. Często w przypadkach fuzji analizy symulacyjne wykorzystywane są przez organy właściwe ds. konkurencji⁵⁶ uzupełniając do powszechnie stosowanych ocen jakościowych.

Analiza symulacyjna odbywa się zwykle w dwóch etapach. W pierwszym tworzy się specyfikację modelu. W tym celu określa się: parametry konkurencji – zazwyczaj konkurencji cenowej i ilościowej, udziały w rynku przedsiębiorstw, koszty krańcowe,

⁵⁶ Patrz np. US-Merger Guidelines, str. 25.

elastyczność cenową i mieszaną popytu, jak również, o ile występuje, wzrost efektywności.⁵⁷ W drugim etapie przeprowadza się właściwą symulację. Bazując na wyliczonych wskaźnikach (jak elastyczność cenowa itp.) i cenach względnie ilościach przed dokonaniem fuzji oferentów, wyliczane są ceny wzgl. ilości po dokonaniu fuzji oferentów.⁵⁸ Na podstawie tych wartości można wyciągnąć wnioski odnośnie efektów cenowych i konsekwencji fuzji dla dobrobytu społecznego.

W zależności od dostępności danych można zastosować różne modele symulacji. Proste modele opierają się jedynie na cenach, wielkości sprzedaży, udziałach w rynku, kosztach krańcowych i elastyczności cenowej popytu. Już na podstawie takiego modelu można uzyskać informacje o zmianach cenowych i zmianach udziału na rynku po dokonaniu fuzji oraz zmiany korzyści społecznych. Przy pomocy tych modeli można dokonywać także analizy wrażliwości, aby zbadać, w jakim stopniu przykładowo reagują ceny na zmiany poszczególnych parametrów. Proste modele mające relatywnie niewielkie wymogi odnośnie danych, bazują jednak najczęściej na restrykcyjnych założeniach, które nie zawsze odpowiadają rzeczywistej sytuacji rynkowej. W takich przypadkach sensowne może być zastosowanie bardziej skomplikowanego modelu. Modele takie wymagają jednak często szerokiego zakresu danych, ale umożliwiają dokładniejszą prognozę. Ogólnie obowiązuje trade off pomiędzy najwyższą z możliwych dokładnością a nakładami potrzebnymi na pozyskanie danych.⁵⁹ W przypadku opinii rzeczoznawczych przedłożonych przez strony należy wziąć pod uwagę, że dysponują one, z powodu braku uprawnień do otrzymywania informacji, dużo mniejszymi możliwościami pozyskiwania danych niż organy antymonopolowe.

8.2 Fuzje wertykalne i konglomeratowe

Często dochodzi do fuzji między przedsiębiorstwami, które nie działają na tym samym rynku, a jednak ich działalność odbywa się na rynkach dostawców lub odbiorców i w ten sposób zawiera się w łańcuchu wytwarzania danego produktu. Przykładowo producent przejmuje dostawcę półproduktu (rynek poprzedzający,

⁵⁷ W tej kwestii pomocna może być często ocena ekonometryczna.

⁵⁸ Wyliczenie opiera się na funkcji zyskowności przedsiębiorstw.

⁵⁹ Szerzej patrz Epstein i Rubinfeld (2004) lub Baker i Rubinfeld (1999).

upstream) lub sprzedawcę swoich produktów (rynek następujący, downstream). Takie fuzje określa się mianem **fuzji pionowej (wertykalnej)**, względnie **integracji pionowej**. W przeciwieństwie do fuzji poziomej fuzje wertykalne nie mają per se skutków ograniczających konkurencję, a w określonych okolicznościach mogą nawet być korzystne dla konkurencji. Fuzje wertykalne mogą jednak wzmocnić pozycję rynkową biorących w nich udział przedsiębiorstw. W takim przypadku efekty unilateralne lub koordynacyjne nie są wywoływane bezpośrednio, ale istnieje możliwość ich pośredniego powstawania. Takie skutki omówione zostaną szczegółowo w rozdziale.

Kiedy fuzja dotyczy przedsiębiorstw, które nie działają ani na tym samym rynku ani na rynkach połączonych wertykalnie, mówi się wtedy o fuzji konglomeratowej. Również takie fuzje mogą wywoływać efekty unilateralne i koordynacyjne. **Efekty unilateralne** mogą wynikać np. z tak zwanych **efektów portfelowych, wzrostu siły finansowej** lub – o ile dane produkty z punktu widzenia popytu lub podaży podlegają relatywnie słabej substytucji – na skutek odejścia **małych** lub **potencjalnych konkurentów**. Mianem efektów portfelowych określa się skutki, które mogą wynikać z rozszerzenia portfela (gamy) produktów. Kiedy przedsiębiorstwo dzięki fuzji rozszerza swój portfel, może ono w określonych okolicznościach wykorzystać swój portfel, aby poprzez **sprzedaż pakietową** wyrugować z rynku konkurentów. Jeżeli przykładowo dwa różne produkty jednego przedsiębiorstwa poszukiwane są przez określoną grupę klientów, to dla przedsiębiorstwa może być opłacalne zmonopolizowanie obu rynków poprzez oferowanie pakietu produktów.⁶⁰ W pewnych okolicznościach opłacalną strategią może być dla przedsiębiorstwa sprzedaż pakietowa w celu powstrzymania potencjalnych konkurentów od wejścia na rynek. Dzięki wiązaniu (i technicznej integracji) produktów oferent może uwiarygodnić fakt, że w przypadku wejścia na rynek produkty nie będą na żadnym rynku sprzedawane osobno, a żeby uniknąć spadku przychodów na wejścia na rynek będzie reagować agresywnymi cenami. Wiązanie produktów służy zatem jako mechanizm odpowiedniego kształtowania bodźców ekonomicznych.

⁶⁰ Szczegóły na ten temat Nalebuff (2004), Greenlee et al. (2004).

W podobny sposób może oddziaływać **wzrost siły finansowej**. Dzięki niemu przedsiębiorstwo biorące udział w fuzji otrzymuje dodatkowe środki finansowe, których może użyć w razie konieczności do wyparcia konkurenta z rynku.⁶¹ Jeśli fuzja konglomeratowi powoduje zniknięcie konkurencji potencjalnej lub brzegowej, to jej skutki porównywalne są do efektów unilateralnych, które zostały przedstawione w ramach fuzji poziomej.

Fuzje konglomeratowe mogą także wywołać **efekty skoordynowane** pomiędzy przedsiębiorstwami biorącymi udział w fuzji a przedsiębiorstwami trzecimi na jednym lub wielu rynkach. Jeżeli poprzez połączenie zwiększa się symetria pomiędzy przedsiębiorstwami, choćby poprzez zrównanie pod względem oferowanego portfela produktów, sprzyja to koordynacji zachowań rynkowych. W szczególnych przypadkach koordynacja może objąć wiele rynków. Faktyczne występowanie efektów skoordynowanych uzależnione jest to od dalszych czynników, opisanych już w punkcie 7.2

9 Unilateralne ograniczenia konkurencji

Celem nadzorowania zjawiska nadużywania pozycji dominującej jest zapobieganie wykorzystywaniu przez przedsiębiorstwa dominujące niepoddanej konkurencyjnej dyscyplinie swobody działania na rynku na niekorzyść osób trzecich, którą umożliwia niekorzystna struktura rynkowa. Przeciwdziałanie nadużywaniu pozycji dominującej powinno w efekcie służyć utrzymaniu otwartego dostępu do rynku innym podmiotom. Równocześnie należy stwierdzić z całą pewnością, że przedsiębiorstwa, które w następstwie dominującej pozycji na rynku faktycznie mogą jednostronnie dyktować warunki kontraktowe, biorą także pod uwagę interesy drugiej strony rynku. To nie przedsiębiorstwa mają być chronione przed konkurencją, ale chroniona ma być konkurencja z uwagi na to, iż przyczynia się do wzrostu wydajności oraz optymalnego zaopatrzenia konsumentów.

⁶¹ Strategia wojny cenowej opłaca się jednak tylko wtedy, jeżeli zysk w jej wyniku uzyskany jest większy niż zysk możliwy do wypracowania bez uciekania się do niej.

9.1 Ograniczanie dostępu do rynku

9.1.1 Ramy analizy

9.1.1.1 POZYCJA NA RYNKU

Pod pojęciem nadużycia, w wyniku którego ograniczony zostaje dostęp konkurentów do rynku należy rozumieć takie praktyki rynkowe, poprzez które możliwości konkurencyjnej działalności innych uczestników rynku, w szczególności rzeczywistych lub potencjalnych konkurentów, ograniczane są „niesprawiedliwie, czy niesłusznie”, w niedozwolony sposób.⁶² Identyfikacja ograniczeń konkurencji nie jest prosta ani łatwa, jako że sukces rynkowy jednego przedsiębiorstwa osiągnięty w warunkach konkurencji ogranicza możliwości rozwoju jego konkurentów. Dopóki sukces ten jest efektem szczególnie wysokiej wydajności nie można podważać go jako niesłusznego, czy niesprawiedliwego ograniczenia. Rodzaje działań, które często traktowane są jako niesprawiedliwe, to zaniżone ceny, określone systemy rabatów, stosunki wyłączności, strategie wiązania towarów lub odmowy dostaw. Dla wielu takich rodzajów zachowania przyjmuje się, że mogą one dopiero wtedy nabierać charakteru niesprawiedliwych utrudnień, gdy równocześnie występuje znaczna dominacja na rynku, która nie pozwala konkurentom na obronę. Z tego powodu określenie dominacji na rynku jest koniecznym warunkiem udowodnienia antykonkurencyjnego nadużywania pozycji dominującej.

9.1.1.2 TEST “AS EFFICIENT COMPETITOR” (RÓWNIE EFEKTYWNEGO KONKURENTA) JAKO PUNKT ODNIESIENIA

Obecnie toczy się dyskusja nad stosowaniem testu **as efficient competitor** (równie efektywnego konkurenta) do identyfikacji rodzajów zachowań dominantów, które ograniczają konkurencję. Myśl przewodnia tego testu jest taka, że chroniona ma być konkurencja, a nie (nieefektywni) konkurenci. Ponieważ nieefektywne przedsiębiorstwa prędzej czy później na skutek braku umiejętności konkurowania z

⁶² Częściowo także **zobowiązania zakupu**, które zobowiązują nabywcę w określonym czasie do zaspakajania całego popytu lub znacznej jego części wyłącznie od jednego oferenta, uznawane są za nadużycie ograniczeń konkurencji i rozważa się wobec nich posunięcia jednostronne. Z uwagi na ich (równoczesny) wielostronny charakter zobowiązania zakupu zostały omówione szczegółowo w poniższym punkcie dot. porozumień pionowych.

efektywnymi przedsiębiorstwami wypadną z rynku, organy ochrony konkurencji, zgodnie z powyższym, powinny chronić tylko takich konkurentów, którzy są przynajmniej tak samo efektywni (*as efficient competitors*), jak przedsiębiorstwa dominujące na rynku. Zgodnie z tym zachowanie przedsiębiorstwa dominującego na rynku należy traktować tylko wtedy jako nadużycie, gdy jest ono w stanie wyprzeć z rynku konkurenta będącego co najmniej tak samo efektywnym.

W istocie test ten ma swoje **ograniczenia i wady**. Nie jest on skuteczny, kiedy przedsiębiorstwo^[u1] produkuje w warunkach wzrastających korzyści skali, tj. średnie koszty zmniejszają się wraz z rosnącą produkcją. Ponieważ przedsiębiorstwo dominujące na rynku dysponuje dużą wielkością produkcji, może ono redukować koszty średnie poprzez rosnące przychody wynikające ze skali tak dalece, że będą one niższe od kosztów średnich przedsiębiorstw wytwarzających mniejszą ilość, nawet wtedy, gdy funkcja kosztów średnich konkurenta przebiegać będzie poniżej funkcji kosztów średnich przedsiębiorstwa dominującego na rynku. W takiej sytuacji mamy do czynienia z typową **zawodnością koordynacji**, która prowadzi do nieefektywnej, ale za to stabilnej równowagi. Ma to szczególne znaczenie w branżach sieciowych, gdzie przedsiębiorstwo o ugruntowanej pozycji w normalnych warunkach posiada z uwagi na swoją wielkość przewagę nad konkurentami.

Ponadto należy wziąć pod uwagę, że także mniej efektywni konkurenci ograniczają możliwości działania przedsiębiorstwa dominującego, próbując zwiększać presję konkurencyjną na danym rynku. Ma to znaczenie przede wszystkim wtedy, gdy dany ^[u2]rynek charakteryzuje się wysokimi barierami wejścia, które utrudniają dostęp potencjalnym konkurentom. Test „równie efektywnego konkurenta” byłby w takim przypadku niewystarczający, gdyby interpretować go w ten sposób, że dopuszcza on eliminowanie z rynku nowych przedsiębiorstw, które z momencie jego zastosowania są co prawda mniej efektywne, ale z czasem mogą być tak samo, a nawet bardziej efektywne niż przedsiębiorstwo dominujące, gdyby dać mu szansę zniwelowania osiągniętej przez dominanta „korzyści pierwszego ruchu” poprzez korzyści wynikające z **efektu uczenia się** i organiczny wzrost.

Ostatecznie głównym zadaniem testu „równie efektywnego konkurenta” jest analiza kosztów przedsiębiorstwa dominującego na rynku. Doświadczenia Federalnego Urzędu Kartelowego w przypadkach sektora energetycznego wskazują na to, że **analiza kosztów** jest metodą skomplikowaną oraz pochłaniającą wiele czasu i

nakładów. Definicja i kalkulacja kosztów są zarówno w teorii jak i w praktyce dyskusyjne i prowadzą często do długotrwałych i niezwykle żmudnych postępowań. Ponadto wiarygodność danych niezbędnych do wyznaczenia wysokości odpowiednich kosztów zależy w dużej mierze od gotowości do współpracy przedsiębiorstwa z organem ochrony konkurencji, w związku z czym pomocne może być wykorzystanie wcześniejszych doświadczeń, na przykład w odniesieniu do benchmarków kosztowych.

Generalnie stosowanie testu as "równie efektywnego konkurenta" może zmniejszyć częstotliwość występowania błędu typu I (zbędne ingerencje), równocześnie jednak z uwagi na ograniczone zasoby urzędu prowadzić może do zwiększenia częstotliwości występowania błędu typu II (zła decyzja o braku ingerencji) na innych rynkach. Zastosowanie testu "równie efektywnego konkurenta" może mieć uzasadnienie w postępowaniach o ogromnym znaczeniu ekonomicznym (np. z uwagi na wielkość i strukturę badanego rynku) wzgl. o wielkim znaczeniu prawnym (precedensy). Oznacza ono jednak poważne **obciążenie przy prowadzeniu postępowania**, gdyby stosować go bezwzględnie w każdym przypadku zachowania jednostronnego.

9.1.1.3 UZASADNIENIE ZACHOWANIA

W dyskusji jako obiektywne powody usprawiedliwiające domniemane zachowanie antykonkurencyjne uznaje się obiektywną konieczność (objective necessity defence), obronę przed konkurencją (meeting competition defence) lub efektywność ekonomiczną danego zachowania (efficiency defence).

Tłumaczenie się **obiektywną koniecznością** (objective necessity defence) można przyjąć wtedy, gdy np. obowiązujące wszystkie przedsiębiorstwa aspekty bezpieczeństwa i ochrony zdrowia uzasadniają badane zachowanie, ponieważ w innym przypadku towary nie mogłyby być produkowane lub sprzedawane.

W odniesieniu do reakcji na **wejście na rynek** nowego konkurenta (meeting competition defence) należy zauważyć, iż powyższa argumentacja obronna stosowana jest regularnie przez przedsiębiorstwa dominujące na rynku, ponieważ ich zachowanie prawie zawsze jest reakcją na zmieniającą się sytuację rynkową. Chociaż każde dominujące na rynku przedsiębiorstwo uprawnione jest bez wątpienia do podążania za własnym interesem ekonomicznym jako swym prawowitym celem, takie zachowanie nie może być dopuszczone, gdy prowadzi do zamknięcia rynku.

Ocena **usprawiedliwienia efektywnościowego** (efficiency defence) wymaga od organu ochrony konkurencji wielkiego nakładu czasu i pracy. Powstaje pytanie, czy takie nakłady w świetle tego, że usprawiedliwienie efektywnościowe udaje się dowieść jedynie w niewielkiej liczbie przypadków, są w ogóle uzasadnione. Generalnie brak jest miejsca na wspomniane usprawiedliwienie przy praktykach drapieżnictwa cenowego i innych nadużyciach cenowych. (por. w tym względzie zapisy w punkcie 9.1.1.2 i dalsze).

9.1.2 Drapieżnictwo cenowe

Dotychczasowa praktyka stosowana wobec drapieżnictwa cenowego opiera się na wyroku Trybunału Europejskiego w sprawie AKZO, według którego ceny zaniżone można scharakteryzować przy pomocy elementu obiektywnego i subiektywnego.

Element obiektywny przedstawia relacje pomiędzy ceną a średnimi kosztami całkowitymi wzgl. średnimi kosztami zmiennymi. **Element subiektywny** zawiera zamiar wyparcia konkurenta z rynku wiążący się ze strategią cenową. Jeżeli cena jest niższa od średnich kosztów zmiennych, wychodzi się z założenia, że istnieje zamiar wyparcia konkurenta z rynku, ponieważ w innym przypadku cena taka byłaby nieracjonalna. W literaturze wskazuje się również na konieczność przeprowadzenia dowodu na możliwość odzyskania poniesionych przez dominanta strat wynikających z zaniżonej ceny, bez której to możliwości trudno jest mówić o stosowaniu cen drapieżnych. Nie istnieje bowiem żadna racjonalna przesłanka dla przedsiębiorstwa do stosowania cen drapieżnych, jeśli nie spodziewa się ono, że poświęcone w czasie wojny cenowej przychody będzie mogło odzyskać w drodze podniesienia cen po wypchnięciu konkurenta z rynku. Jeżeli oczekiwanie takie opiera się na nieracjonalnej ocenie sytuacji przez przedsiębiorstwo, wskazuje to raczej, iż jest ono nieefektywne lub posiada niewielką siłę rynkową, w związku z czym nie ma podstaw do interwencji ze strony organu ochrony konkurencji. Z tego powodu można uznać dowód na możliwość odzyskania poświęconych w wojnie cenowej przychodów za bardzo ważny dla oceny danego przypadku.⁶³

⁶³ Niezależnie od tego kalkulacja możliwości odzyskania utraconych przychodów może być bardzo trudna lub wręcz niemożliwa, jeżeli np. stosuje się strategię cen drapieżnych w tym celu, aby stworzyć odpowiednio znaczącą reputację, w celu zapobiegania wejściom konkurentów także na inne rynki.

W przypadku cen, które mieszczą się w przedziale między średnimi kosztami zmiennymi a średnimi kosztami całkowitymi, należy udowodnić zamiar wyparcia z rynku poprzez dalsze dowody, ponieważ taka cena przynajmniej krótkoterminowo może być racjonalna.

Jakkolwiek ceny powyżej średnich kosztów całkowitych są w przypadkach wyjątkowych – np. na skutek silnych korzyści skali – cenami zaniżonymi, zazwyczaj nie traktuje się ich w ten sposób.

Test AKZO można również zmodyfikować. Zasadne wydaje się potraktowanie jako punktu odniesienia średnich kosztów możliwych do uniknięcia (zamiast średnich kosztów zmiennych). Poniżej tego progu można podejrzewać zamiar wypchnięcia konkurentów z rynku. W przypadku cen pomiędzy średnimi kosztami możliwymi do uniknięcia a średnimi kosztami całkowitymi dowód zamiaru wyparcia konkurenta z rynku – choćby poprzez wykorzystanie siły finansowej lub w celu stworzenia reputacji na innych rynkach – wciąż jest możliwy.

9.1.3 Systemy rabatowe

W dotychczasowej praktyce Komisji i orzecznictwie Europejskiego Trybunału Sprawiedliwości w kwestii nieprawidłowości w systemach rabatowych rozróżnić należy pomiędzy niestanowiącymi z reguły problemu rabatami ilościowymi, a z natury rzeczy problematycznymi rabatami lojalnościowymi i celowymi. Nadużycia w zakresie systemów rabatowych wiążą się zasadniczo z kwestią, czy poprzez udzielenie rabatu powstaje efekt „zasysania popytu” lub czy rabatowi odpowiada osiągnięta korzyść ekonomiczna, która przenoszona jest dalej na stronę popytu. W zasadzie kwestia oceny, czy doszło do nadużycia, zorientowana jest prawie wyłącznie na stronę kosztową, tj. wszelkie inne uzasadnienia dla udzielenia rabatu nie są brane pod uwagę. Takie podejście z ekonomicznego punktu widzenia może być często nieodpowiednie w odniesieniu do konkretnych warunków gospodarczych.

Bardziej merytoryczny jest podział na warunkowe i bezwarunkowe systemy rabatów. Bezwarunkowe rabaty udzielane są niezależnie od zachowania po stronie popytu i charakteryzują się niższymi cenami. Nie są one w większości przypadków kwestionowane z punktu widzenia prawa konkurencji. W każdym przypadku selektywne udzielanie rabatów bezwarunkowych może być uznane za przejaw zamiaru wypierania z rynku.

Rabaty warunkowe mogą z jednej strony mieć charakter przyrostowy, tzn. po przekroczeniu określonego progu obrotów udzielany jest wyższy upust w odniesieniu do dalszych zamawianych jednostek. Niebezpieczeństwo negatywnego oddziaływania rabatów przyrostowych na konkurencję należy generalnie postrzegać jako niewielkie. Należy je oceniać jako nadużycie tylko wtedy, gdy cena po udzieleniu rabatu może być uznana za cenę drapieżną (zaniżoną) (por. 9.1.2), tzn. gdy jest ona niższa od średnich kosztów całkowitych przedsiębiorstwa dominującego na rynku. Selektywne przyznawanie rabatów może być w tej sytuacji uznane za przejaw strategii wypierania z rynku.

Generalnie jako problematyczne należy oceniać takie rabaty warunkowe, przy których od przekroczenia określonego progu obrotów rabatowi podlegają wszystkie dotychczas zamówione jednostki (tzw. „**rabaty retroaktywne**”). Mogą one pociągać za sobą znaczne koszty zmiany (switching costs), jeżeli udzielony rabat jest odpowiednio wysoki. W przypadku rabatów retroaktywnych kwestia nadużycia jest uzależniona od efektu „ssania” powstałego na skutek udzielonego rabatu. Jeżeli poprzez rabat udział w rynku, który musi osiągnąć konkurent, aby odnosić korzyści przy takiej samej strukturze kosztów i dostępnym systemie rabatowym, jest większy od spornego udziału na rynku⁶⁴, wtedy system rabatowy szkodzi konkurencji i to niezależnie od tego, czy w konsekwencji stosowane są ceny zaniżone. W takiej sytuacji efektywni konkurenci zostaliby wyparci z rynku a konkurencja byłaby tym słabsza, im bardziej utrudniony stałby się dostęp do rynku (por. ww. rozdział 7.1.5).

Stwierdzenie to nie obowiązuje w drugą stronę w takim sensie, że nie można jednoznacznie stwierdzić, iż rabaty retroaktywne nie stanowią nadużycia, jeśli udział w rynku niezbędny konkurentowi do prowadzenia zyskowej produkcji jest równy lub niższy od wielkości popytu, która pozostaje na rynku po uwzględnieniu sprzedaży dominanta. W tej sytuacji wynikający z rabatu efekt „ssania” musi być oceniany ex ante. Stosunek pomiędzy rzeczywiście dokonanymi zakupami a udziałem w rynku

⁶⁴ Oznacza on udział w ogólnym popycie, który nie „musi” być nabyty u dominującego na rynku przedsiębiorstwa. „Konieczność” nabywania części towarów u dominanta rynkowego może przykładowo być konsekwencją korzyści wynikających z różnorodności produktu lub reputacji zastosowanie tego dobra w produkcji, względnie faktu, iż wprowadzenia go do oferty handlowej oczekuje wielu nabywców. Podobny skutek może wynikać z ograniczenia mocy produkcyjnych konkurentów, niepodzielności lub nieelastycznego popytu.

jest natomiast oceniany ex post, co nie odzwierciedla bodźców ekonomicznych przedsiębiorstwa w momencie podejmowania decyzji o zakupie.

9.1.4 Sprzedaż wiązana i pakietowa

Sprzedaż wiązana ma miejsce wtedy, gdy sprzedaż produktu podstawowego możliwa jest tylko wraz z produktem powiązanym, tj. produkt podstawowy nie jest dostępny osobno. Sprzedaż pakietowa ma miejsce wtedy, gdy oferowany jest pakiet wielu produktów, przy czym odróżnia się "**pure bundling**" (pojedyncze produkty nie są oferowane) oraz "**mixed bundling**" (pojedyncze produkty są oferowane osobno, jednak na pakiet produktów udzielany jest rabat).

Sprzedaż wiązana i pakietowa nie prowadzi co prawda często do ograniczenia konkurencji, może być ona jednak mimo wszystko stosowana przez przedsiębiorstwa w celu zamknięcia konkurentom dostępu do rynku, różnicowania cen lub w celu osiągnięcia wyższych cen. W szczególności sprzedaż wiązana umożliwia wykorzystanie pozycji dominującej na rynku produktu podstawowego celem poprawienia pozycji rynkowej na rynku produktu związanego (**efekt dźwigni**).

Udowodnienie nadużycia wymaga stwierdzenia, oprócz pozycji dominującej na rynku produktu podstawowego (opanowanie rynku produktów powiązanych nie jest koniecznym wymogiem), że rynek produktu podstawowego i rynek produktu z nim powiązanego stanowią **osobne rynki**, oraz że poprzez wiązanie produktów konkurencja została ograniczona lub osłabiona i że nie ma uzasadnienia dla takiego powiązania.

Aby ukazać ograniczenie konkurencji, należy w pierwszym etapie zidentyfikować takich **konsumentów (grupy konsumentów)**, którzy na skutek sprzedaży związanej lub pakietowej związani są z dominatem na rynku i z tego powodu nie wchodzi w rachubę jako nabywcy produktów konkurentów. Podczas gdy nabywcy w przypadku sprzedaży związanej i pure bundling w sposób oczywisty związani są z dominatem na rynku, przy mixed bundling należy dokładniej zbadać **koszty przyrostowe**⁶⁵, aby wykazać, że rywale nie mają możliwości konkurowania o nabywców.

⁶⁵ Koncepcja kosztów przyrostowych przypomina koncepcję kosztów możliwych do uniknięcia (por. rozdział 3.2.2), przy czym rozpatruje się tu nie zmniejszenie produkcji (unikanie ostatniej jednostki), lecz jej poszerzenie. Koszty przyrostowe obejmują obok kosztów zmiennych także

W drugim etapie należy udowodnić w ramach całościowej oceny, że **rynek został w pełni zamknięty dla konkurentów**. Należy wykazać, że zidentyfikowani w pierwszym etapie konsumenci stanowią znaczącą większość na badanym rynku. Efekty sieciowe i skali jak również bariery dostępu do rynku stanowią ważne aspekty tego badania. W każdym razie należy uwzględnić, że zamknięcie rynku jest procesem dynamicznym, w którego ramach poprzez sprzedaż wiązaną i pakietową pozycja dominująca na rynku produktu głównego przenoszona jest na rynek produktu powiązanego (efekt dźwigni). Stąd możliwość przenoszenia dominacji rynkowej jest istotna dla udowodnienia ograniczenia dostępu do rynku i to nawet, w razie konieczności, niezależnie od badanego udziału na rynku.

Podczas analizy mogą wystąpić problemy, o ile **kalkulacja przyrostowej ceny** jakiegoś pojedynczego produktu sprzedawanego w pakiecie okazuje się skomplikowana i wymaga wielkiej liczby danych, szczególnie wtedy, gdy na pakiet produktów udzielane są rabaty lub gdy nabywcy dokonują zakupów w różny sposób. Pomocniczo można założyć, że jeden z konkurentów rzeczywiście został wyłączony lub wyparty z rynku.

9.1.5 Odmowa dostaw

Punktem wyjścia jest tu swoboda zawierania umów przez podmioty gospodarcze (zasada swobodnej wymiany; por. rozdział 2.1), i to niezależnie od tego, czy posiadają one mocną pozycję rynkową czy nie. W określonych okolicznościach odmowa dostawy stanowi jednak faktycznie nadużycie posiadanej siły rynkowej. Toma to miejsce zwłaszcza wtedy, gdy jest ona stosowana jako środek do osiągnięcia innych niezgodnych z prawem celów (np. umowa wyłączności i sprzedaż wiązana). W ramach analizy należy odróżnić potencjalnie szkodliwe skutki utrudnienia dostaw w przypadku rynków, na których dominant jest odbiorcą i takich, an których jest dostawcą (**upstream-downstream**). Najwięcej problemów można dostrzec w przypadkach odmowy dostaw na tych drugich rynkach.

Należy wskazać przede wszystkim dwie potencjalne struktury nadużycia:

wszelkie specyficzne dla danej produkcji koszty stałe a nie jedynie część kosztów stałych możliwą do uniknięcia. W tym kontekście długookresowe średnie koszty przyrostowe są zawsze większe niż średnie koszty możliwe do uniknięcia.

1. **Odmowa** dostawy **półproduktu** przez przedsiębiorstwo dominujące na rynku, o ile półprodukt jest niezbędny do produkcji, o ile odmowa dostawy nie znajduje obiektywnego uzasadnienia i o ile przypuszczalnie spowoduje to negatywne skutki dla konkurencji. W przypadku **odmowy licencji** należy do tych warunków dodać przeszkody w rozwoju nowych produktów lub usług.
2. **Zerwanie** istniejącego stosunku dostaw przez przedsiębiorstwo dominujące na rynku, o ile pociąga to za sobą prawdopodobny efekt ograniczenia dostępu osłabiający konkurencję rynkową i nie jest ono uzasadnione ani obiektywnie, ani poprzez zwiększenie efektywności ekonomicznej;

Orzeczenie Europejskiego Trybunału Sprawiedliwości w sprawie Bronner podkreśla, że odmowa dostawy musi być **konieczna („indispensable“)**, przy czym tej konieczności stawiane są bardzo wysokie wymagania: towar nie może być ani zastępowalny ani możliwy do wytworzenia we własnym zakresie. Jest to zgodne z powszechną definicją essential-facilities (niezbędnych urządzeń). Kryterium konieczności nie jest jednak miarodajne dla grupy przypadków zerwania istniejącego stosunku dostaw, do których można zaliczyć także przypadki „ściskania marż” (margin squeeze). Mówi się o nich, gdy spółki córki zajmujące się sprzedażą przedsiębiorstwa zintegrowanego pionowo oferują na rynku towar po niższych cenach niż te, które nalicza się przedsiębiorstwom trzecim za zakup tego samego dobra.

Za **obiektywne powody uzasadniające** odmowę dostaw przedsiębiorstwa dominującego na rynku mogą zostać uznane: brak zdolności płatniczej lub wiarygodności kredytowej nabywcy, a także dążenie do pokrycia nakładów na wielkie i obciążone wysokim ryzykiem inwestycje. Ograniczone czasowo wyłączenie z dostępu do jednego z niezbędnych dla konkurenta czynników produkcji może być dopuszczalne, aby utrzymać na rynku silne bodźce do inwestycji i innowacji. W odniesieniu do bodźców do innowacji istnieją tak czy inaczej uregulowania chroniące własność intelektualną. W odniesieniu do bodźców inwestycyjnych należy zbadać, mając na uwadze odpowiednią rekompensatę pieniężną jako alternatywę, czy (w przypadku essential facilities) czasowe dopuszczenie odmowy dostaw jest najłagodniejszym środkiem. W wyniku takiego działania bowiem aktualni lub potencjalni konkurenci mogą wypaść z rynku. Jeśli ceny – zgodnie z oczekiwaniami –

skalkulowane zostaną w sposób rzetelny i z uwzględnieniem poniesionych nakładów inwestycyjnych, to dominującemu przedsiębiorstwu powinno być z ekonomicznego punktu widzenia obojętne, czy dostarczy dany produkt, czy też nie.

9.1.6 Rynki wtórne

Mianem rynków wtórnych określa się takie rynki, na których towary będące przedmiotem obrotu użytkuje się wyłącznie w połączeniu z zakupionymi wcześniej innymi dobrami komplementarnymi (np. zszywacz i zszywki, sprzęt i oprogramowanie komputerowe oraz usługi wspomagające).

Na rynkach wtórnych ograniczające konkurencję nadużywanie pozycji dominującej może wystąpić wtedy, gdy rynki będą musiały zostać wyznaczone w sposób zawężony, to znaczy wyłącznie w odniesieniu do produktów komplementarnych producenta dóbr pierwotnych, dlatego że nie istnieją substytuty innych producentów wyrobów komplementarnych lub gdy koszty zamiany na poziomie produktów pierwotnych są wysokie. To z kolei jest prawdopodobne, gdy przychody ze sprzedaży towarów używanych są względnie niewielkie lub wymagane są duże kolejne inwestycje (np. przeszkolenie znacznej części personelu). Wtedy siła rynkowa na rynku wtórnym jest prawdopodobna, ponieważ istnieją na nim zwykle techniczno-prawne (patenty) lub strategiczne bariery dostępu.⁶⁶

Nadużycie wystąpiłoby, gdyby przedsiębiorstwa przeszkadzały potencjalnym konkurentom (którzy jednak nie konkurują na rynku dóbr pierwotnych) w wejściu na rynek wtórny. Mogłoby to się odbywać poprzez np. sprzedaż wiążaną lub odmowę dostaw. Ponadto możliwe byłyby utrudnienia uwarunkowane subwencjonowaniem skośnym na (ew. podlegającym silniejszej konkurencji) rynku towarów pierwotnych dla konkurentów, którzy nie działają na rynku wtórnym i dla których nie otwierają się takie możliwości subwencjonowania skośnego. W odniesieniu do sprzedaży pakietowej analiza, czy na rynku pierwotnym stosowane są ceny zaniżone, byłaby prawie niemożliwa do wykonania (por. 9.1.4).

⁶⁶ Ponieważ "wet za wet", czyli działania odwetowe przy wejściu jednego z konkurentów na rynku pierwotnym na jego własny rynek wtórny, jest strategią dominującą (por. rozdział 4.5.2).

9.2 Dyskryminacja cenowa

Armstrong, M.; Vickers, J. (1993): Price Discrimination, Competition and Regulation. In: Journal of Industrial Economics, Vol. 41, str. 335-360. Chen, Y. (1999): Oligopoly Price Discrimination and Resale Price Maintenance. In: RAND Journal of Economics, Vol. 30, str. 441-455. Corts, K. (1998): Third-Degree Price Discrimination in Oligopoly. In: RAND Journal of Economics, Vol. 29, str. 306-323. Konkurrensverket (Ed.) (2005): The Pros and Cons of Price Discrimination. Stockholm. Nilssen, T. (1992): Two Kinds of Consumer Switching Costs. In: RAND Journal of Economics, Vol. 23, str. 579-589.

Dyskryminacja cenowa (lub różnicowanie cen) oznacza, że pewne dobro sprzedawane jest po różnych cenach a **różnice cenowe** nie oddają odpowiednich różnic kosztowych. Motywacja do różnicowania cen przez oferenta bierze się stąd, że oferując każdemu nabywcy inną cenę na określony towar, w miarę możliwości stosownie do jego indywidualnej gotowości płatniczej, przejmuje nadwyżkę konsumenta.

Dyskryminacja cenowa zakłada, że **gotowość do płacenia** różnych nabywców da się obliczyć bez większych nakładów w stopniu wystarczająco dokładnym. Następnie musi zaistnieć możliwie skuteczne ograniczenie uprawnień do odsprzedaży określonego towaru, ponieważ w innym wypadku dochodziłoby do transakcji arbitrażowych⁶⁷, przeciwdziałających możliwościom dyskryminacji.

Można wyróżnić trzy rodzaje dyskryminacji cenowej. **Zróżnicowanie cen pierwszego stopnia** występuje wtedy, gdy każdy nabywca musi zapłacić cenę odpowiednio do indywidualnej gotowości płatniczej. Z uwagi na związane z tym wysokie koszty informacyjne, rzadko można zaobserwować tę formę dyskryminacji cenowej w praktyce. O wiele prostsza do zrealizowania jest **dyskryminacja cenowa drugiego stopnia**, która jest równoznaczna z rabatem ilościowym. Cena za jednostkę jest tu uzależniona od liczby jednostek kupowanych. Ten rodzaj dyskryminacji cenowej znajduje często odbicie w taryfach dwuczęściowych, które składają się przykładowo z opłaty podstawowej niezależnej od ilości oraz komponentów uzależnionych od poziomu zakupów. **Zróżnicowanie cenowe trzeciego stopnia** występuje wtedy, gdy można wyróżnić grupy nabywców o różnej

⁶⁷ Z arbitrażem mamy do czynienia wówczas, gdy powstające w tym samym czasie różnice pomiędzy cenami danego dobra wykorzystuje się do uzyskania dochodu poprzez kupowanie go po niższej cenie i natychmiastowe odsprzedawanie po cenie wyższej.

elastyczności popytu. Wtedy pobierana jest od każdej grupy nabywców dopasowana do niej cena Cournota (por. rozdział 4.3).

Jeżeli dla przedsiębiorstwa dominującego na rynku, ustalającego ceny Cournota istnieje możliwość dyskryminacji cenowej pierwszego lub drugiego stopnia, to korzystne jest dla tego przedsiębiorstwa oferowanie ilości większych niż w rozwiązaniu Cournota, a **wielkość sprzedaży** jest taka sama jak w przypadku konkurencji doskonałej. Unika się w ten sposób szkodliwej dla gospodarki narodowej straty społecznej, tj. przypadku, w którym konsumenci są skłonni zapłacić cenę wyższą od kosztów krańcowych, jednak nie są gotowi na zapłacenie ceny monopolowej.

Nabywcy znajdują się jednak tutaj w gorszej sytuacji niż przy wyniku monopolowym, ponieważ tracą **nadwyżkę konsumenta** w całości (w przypadku dyskryminacji cenowej pierwszego stopnia) lub częściowo (w przypadku dyskryminacji cenowej drugiego stopnia) na rzecz dominanta. Nawet jeżeli można liczyć się z tendencją do optymalnego z punktu widzenia gospodarki narodowej zaopatrzenia w dobra przy różnicowaniu cen, to mimo to wskazane są ingerencje organów ochrony konkurencji i to niezależnie od opisanych powyżej skutków dla podziału nadwyżki, które mogłyby być z normatywnego punktu widzenia niepożądane. Podmiot dominujący mógłby bowiem wykorzystać dodatkowe zyski, aby swoją pozycję rynkową rozszerzyć na rynki sąsiadujące. Możliwe negatywne konsekwencje dyskryminacji cenowej drugiego stopnia dla konkurencji zostały opisane w związku z systemami rabatowymi w poprzednim rozdziale (ograniczanie dostępu do rynku).

Niezależnie od tego rośnie niebezpieczeństwo **spadku dobrobytu w aspekcie dynamicznym**, ponieważ obie funkcje konkurencji – dopasowywanie się do zmieniających się warunków ramowych i wspieranie postępu technicznego – nie działają tu optymalnie. Powodem jest występowanie impulsów do kierowania aktywności innowacyjnej na określone dziedziny i zaniedbywania innych dziedzin, takich jak redukowanie, względnie przełamanie subaddytywności w funkcji kosztów, ponieważ sukces takich działań innowacyjnych umożliwiłby wejście na rynek nowych konkurentów. W ten sposób można by wyjaśnić, dlaczego zakłady energetyczne kierują swoje poczynania innowacyjne na wielkie technologie a zaniedbują w znacznym stopniu rozwój zdecentralizowanych technologii wytwarzania energii.

9.3 Zachowania eksploatacyjne

Jak przedstawiono w punkcie 4.3 straty dobrobytu powstają przez to, że dominujące na rynku przedsiębiorstwo ustala cenę, która jest wyższa od kosztów krańcowych. Narzut cenowy i związana z nim strata dobrobytu jest tym większa, im większa jest siła rynkowa przedsiębiorstwa. Zgodnie z art. 82 Traktatu Wspólnot Europejskich pobieranie zawyżonych cen stanowi **nadużywanie siły rynkowej**.

Korzystanie z siły rynkowej poprzez zawyżanie cen może oddziaływać na różne rynki. Z początku przedsiębiorstwo na opanowanym rynku z uwagi na swoją dominującą pozycję może samo żądać ceny, która znacznie przewyższa koszty krańcowe. Takie zachowanie jest szczególnie racjonalne w przypadkach, gdy dominujące przedsiębiorstwo nie musi się liczyć z nowymi wejściami na rynek. Ma to praktyczne znaczenie szczególnie na rynkach, na których funkcjonują **nieefektywne regulacje** zwłaszcza w odniesieniu do **essential facilities**. Przykładem może być transport lotniczy pomiędzy dwoma znaczącymi lotniskami tranzytowymi, tzw. hubs. Liczba uprawnień do startów i lądowań (tzw. slots) jest tu bardzo ograniczona i na skutek ingerencji państwa z reguły ogranicza się do dawnych państwowych towarzystw lotniczych. Ponieważ sloty – o ile są używane – prawdopodobnie w sposób nieefektywny pozostają w rękach właścicieli uprawnień, nie muszą oni się obawiać wchodzenia nowych konkurentów i mogą na określonych trasach długoterminowo stosować zawyżone ceny.

Oprócz tego istnieje również możliwość, że ekonomiczna eksploatacja odbiorców ma miejsce na **rynku wtórnym**, którego produkty są uzupełnieniem produktów sprzedawanych na zdominowanym rynku (por. rozdział 5.3). Aby wyjaśnić kwestię, czy jedno z przedsiębiorstw dominujących na rynku wtórnym dopuszcza się nadużycia polegającego na stosowaniu wygórowanych cen^[u3], należy uwzględnić również konkurencję na rynku pierwotnym. Przy funkcjonującej konkurencji na rynku pierwotnym cena za dobra komplementarne jest istotnym parametrem konkurencji, ponieważ konsumenci podejmując decyzję o zakupie porównują ceny pakietu towarów (np. kupując drukarkę uwzględniają również cenę wkładu, a kupując samochód - zużycie benzyny i koszty części zamiennych). Jeżeli na rynku pierwotnym konkurencja jest niewielka, możliwe jest stosowanie cen zawyżonych na rynku wtórnym.

Nawet jeżeli rynek pierwotny nie jest zdominowany a ceny pakietu są z początku konkurencyjne, istnieje możliwość zmiany zachowania (**hold up**; por. rozdział 5.3), która powoduje, że właściciele dobra pierwotnego na skutek zmiany polityki transakcyjnej i cenowej będą wyzyskiwani na rynku wtórnym po wprowadzeniu tam podwyżki cen. Czy będzie to możliwe zależy od tego, jak długi okres uwzględniają w swoich obliczeniach konsumenci i jak wysoki jest poziom dostępnych informacji o dalszych kosztach związanych z użytkowaniem dobra podstawowego oraz jak wysokie byłyby nieodwracalne koszty zmiany dostawcy dóbr pierwotnych. Poza tym zachowanie eksploatacyjne w stosunku do konsumentów oznacza pozbywanie się reputacji, co ma negatywny wpływ na gotowość zakupu produktów oferenta przez nowych nabywców. Takie zachowanie jest jednoznacznie racjonalne tylko wówczas, gdy przedsiębiorstwo planuje wyjście z rynku.

Ściganie zawyżonych cen jako nadużycia eksploatacyjnego wiąże się z pewnymi problemami i z tego powodu bywa przedmiotem krytyki.⁶⁸ Najpierw należy stwierdzić, czy ceny są rzeczywiście zawyżone. Za zawyżone uznaje się ceny znacznie przekraczające koszty krańcowe. Do tego celu potrzebny jest szeroki zakres informacji o kosztach krańcowych przedsiębiorstwa. Pomiędzy organem ochrony konkurencji a badanym przedsiębiorstwem zachodzi jednak poważna **asymetria informacyjna**. Informacje o indywidualnych kosztach krańcowych dostępne są wyłącznie w badanym przedsiębiorstwie, na którego informacje zdany jest organ ochrony konkurencji badający koszty krańcowe. W każdym razie przedsiębiorstwo posiada motywację, aby podawać wyższe koszty krańcowe niż te poniesione w rzeczywistości, aby w ten sposób prezentować wysoką cenę jako niewiążącą się z eksploatacją. Ponadto organ ochrony konkurencji staje przed problemem oceny, które składniki kosztów należy przyporządkować procesowi produkcji.

W celu stwierdzenia zawyżenia cen w stopniu nadużycia można w codziennej pracy posłużyć się metodami alternatywnymi wobec analizy kosztów. Metody te zostały przedstawione w rozdziale 10.1.2.3. Służą one do wyliczania hipotetycznej ceny konkurencyjnej szczególnie na podstawie danych z rynków porównywalnych.

⁶⁸ Analiza krytyczna i kilka konkretnych przykładów patrz Evans i Padilla (2004).

Generalny problem przy ocenie marży cenowej i kosztów krańcowych polega na tym, że uwzględnienie kosztów stałych przy konstruowaniu nowych produktów lub tworzeniu infrastruktury nie jest brane pod uwagę. W wielu gałęziach przemysłu takie koszty stałe są znaczące w porównaniu do kosztów krańcowych – jest tak np. w takich branżach, jak np. produkcja oprogramowania, wytwarzanie farmaceutyków i telekomunikacja. Przedsiębiorstwa zainwestują w rozwój nowych produktów tylko wtedy, tj. poniosą koszty stałe tylko wtedy, jeśli będą mogły w krótkim czasie odzyskać te nakłady w formie odpowiednich zysków jednostkowych – a więc stosując ceny wyższe od kosztów krańcowych. Aby to sobie zapewnić, przedsiębiorstwo może przykładowo zgłosić wniosek patentowy na nowo stworzony produkt, co pozwala mu przez określony czas utrzymać pozycję monopolisty w odniesieniu do tego produktu. Jeżeli na skutek zakazu stosowania wyższych cen przedsiębiorstwo nie mogłoby korzystać z przywilejów patentowych, to straciłoby impuls do rozwoju wzgl. wprowadzania nowych produktów. W tym względzie występuje **trade off pomiędzy statycznymi a dynamicznymi funkcjami konkurencji**, który należy także wziąć pod uwagę przy badaniu cen zawyżonych.

10 Porozumienia wielostronne

10.1 Porozumienia poziome / horyzontalne

Porozumienia poziome pomiędzy konkurentami dotyczą produktów podobnych wzgl. zamiennych. Dlatego też przedsiębiorstwa mogą być zainteresowane dążeniem do osiągnięcia kooperatywnego wyniku rynkowego i podziałem między siebie osiągniętych w ten sposób wspólnie korzyści. (por. rozdział 7.2). Np. małe przedsiębiorstwo może skorzystać z tego, że konkurent mający silną pozycję rynkową podniósł swoją cenę powyżej poziomu konkurencyjnego. Podniesienie ceny powyżej poziomu panującego w warunkach konkurencji zwiększa wprawdzie nadwyżkę producenta, ale zmniejsza nadwyżkę ogólną, a więc i poziom dobrobytu. Układ bodźców ekonomicznych w stosunkach pomiędzy konkurentami sprzyja zatem szkodliwej dla konkurencji koordynacji zachowań (np. w postaci zмовы cenowej).

O ile w ramach kolektywnej dominacji na rynku dzieje się to bez szczegółowych porozumień pomiędzy uczestnikami, istnieje także możliwość, że przedsiębiorstwa

porozumieją się szczegółowo co do działania kooperatywnego. Takie porozumienie może obejmować ustalenia co do cen lub ich komponentów, wielkości zbytu, podziału rynku geograficznego lub produktowego itp. Z uwagi na szkodliwe dla konkurencji skutki porozumienia pomiędzy konkurentami zgodnie z art. 81 TWE są zasadniczo zakazane.

Tak jak „milcząca koordynacja (tacit collusion), również szczegółowe porozumienia nie zawsze są stabilne. W przypadku tajnych porozumień uczestnicy nie mają możliwości występowania na drodze prawnej przeciwko łamaniu powziętych ustaleń. Takie porozumienia muszą być stabilne same z siebie. Oznacza to, że uczestnik nie może odnosić korzyści, kiedy odstępuje od porozumienia. Stabilizacji porozumień sprzyjają ponadto opisane w powyżej (rozdział 7.2) czynniki: przejrzystość rynku, brak zewnętrznej konkurencji, mała asymetria uczestników, mało elastyczny popyt itp. W przeciwieństwie do milczącej koordynacji, w przypadku szczegółowych porozumień łatwiej jest osiągnąć wynik kooperacyjny i utrzymywać go. Przedsiębiorstwa mogą porozumieć się w bezpośrednich negocjacjach co do wspólnej ceny i nie muszą znajdować kompromisów metodą prób i błędów lub poprzez kosztowne sygnalizowanie. Poza tym prościej jest reagować na zmiany warunków rynkowych w bezpośredniej komunikacji.

10.1.1 Dopuszczalne porozumienia

Nie każde porozumienie pomiędzy przedsiębiorstwami prowadzi do obniżania dobrobytu społecznego. Porozumienia w określonych warunkach mogą mieć także pozytywne skutki. Dlatego też porozumienia między przedsiębiorstwami mogą być dopuszczone na mocy art. 81 ust. 3 TWE. Dobrobyt społeczny mogą podnosić porozumienia w sprawie rozwoju nowych produktów. Pojedyncze przedsiębiorstwa nie są w stanie samodzielnie rozwijać i produkować określonych wyrobów. Zwykle w takich przypadkach trzeba ponieść znaczne nakłady kosztów stałych na rozwój lub produkcję, których nie może odzyskać ani jedno przedsiębiorstwo samodzielnie, ani kilka konkurujących ze sobą przedsiębiorstw, dlatego bez współpracy konkurentów niektóre produkty w ogóle by nie powstały. Ponadto w przypadku projektów badawczych mogą powstawać efekty typu spillover, co powoduje redukcję indywidualnego zainteresowania takimi projektami. Spillover oznacza, że wyniki projektu – przykładowo tańsza technologia produkcji – uzyskują nieodpłatnie także

inne firmy, a przedsiębiorstwo, które projekt przeprowadziło, nie może temu przeszkodzić.⁶⁹ Równocześnie z punktu widzenia dobrobytu pożądane jest, aby innowacja znalazła jak najszersze zastosowanie.

Kooperacja zwiększająca dobrobyt w przypadku badań i rozwoju odbywa się często w postaci Joint Ventures. Istnieje niebezpieczeństwo, że kooperacja nie ograniczy się do badań i rozwoju, lecz doprowadzi do wspólnych zachowań uczestniczących w niej przedsiębiorstw na rynkach produktowych. Niebezpieczeństwo jest tym większe, im większa jest siła rynkowa uczestniczących przedsiębiorstw na poszczególnych rynkach produktowych oraz im bardziej aktywność Joint Ventures skierowana jest nie tylko na rozwój, ale także na produkcję i dostarczanie produktów na rynek. Działalność Joint Ventures powinna być z tego powodu ograniczona do określonych dziedzin. Kooperacja pomiędzy konkurentami może być również wynikać z porozumień dotyczących wzajemnego korzystania z patentów lub ze wspólnego przyjęcia określonych standardów produktów.⁷⁰ Również takie uzgodnienia kryją jednak w sobie niebezpieczeństwo sprzyjania koordynacji zachowań rynkowych ich uczestników.

Ponadto porozumienia poziome pomiędzy określonymi przedsiębiorstwami mogą zwiększać dobrobyt społeczny, jeżeli np. ograniczają siłę ekonomiczną drugiej strony rynku. Jako przykład na to można wymienić organizacje handlowe lub związki kupieckie.

10.1.2 Niedopuszczalne porozumienia („kartele hardcore“)

W przypadku porozumień dotyczących ceny, wielkości produkcji, terytoriów sprzedaży lub udziału w rynku (tzw. kartele hardcore) należy z reguły wykluczyć efekty przynoszące wzrost dobrobytu. Dlatego też w wielu systemach ochrony konkurencji stanowią one tzw. twarde ograniczenia konkurencji i nie mogą zostać dopuszczone. Dla wykrycia takich karteli, istotne jest zrozumienie ich istoty,

⁶⁹ Wiedza zdobyta dzięki takiemu projektowi przepływa do innych przedsiębiorstw, np. jeżeli wyniki jeszcze nie są chronione patentem i mogą być powielane przez konkurentów. Sprzyja temu np. przepływ pracowników między przedsiębiorstwami. Por. również opis czynników zewnętrznych i zawodności rynku w rozdziale 5.1.

⁷⁰ Szczegółową analizę wpływu porozumień dotyczących ochrony patentowej i standaryzacji produktów na dobrobyt można znaleźć w: Motta (2004), str. 207 i dalsze oraz podana tam literatura.

pozwalające na zgromadzenie danych wskazujących na możliwość występowania kartelu.

10.1.2.1 ISTOTA KARTELI HARDCORE

Kartele to formy organizacyjne i porozumienia kontraktowe pomiędzy zakładami, przedsiębiorstwami lub państwami. Kartel z punktu widzenia polityki cenowej i produkcyjnej porównywalny jest z monopolem złożonym z wielu zakładów. Powstaje wtedy, gdy firmy w jednej branży łączą się ze sobą celem koordynacji wielkości produkcji lub działań dotyczących zbytu. W odróżnieniu od monopolu kartel nie ma możliwości zamykania zakładów, ponieważ do niego nie należą. Raczej chodzi tu o samodzielne przedsiębiorstwa, które nie zgodziłyby się na swoje zamknięcie, ponieważ nie mogą liczyć, że inni członkowie kartelu długoterminowo zapewnią im udziały w swoich zyskach. (por. rozważania dot. niepewności w rozdziale 5.3.2 i teorii gier w rozdziale 5.3.3)

Maksymalizując zysk **produkcję kartelu** można uzyskać, jak w przypadku monopolu (por. rozdział 4.3), poprzez zrównanie kosztów krańcowych każdego członka kartelu z przychodem krańcowym dla całego rynku. W ten sposób suma zysków jest maksymalizowana dla wszystkich uczestników kartelu. W porównaniu z rynkami konkurencyjnymi zmniejsza się wielkość podaży a cena idzie w górę.

Wraz z rosnącą liczbą członków kartelu spada zarówno produkcja każdej firmy, jak i cena rynkowa. W ten sposób maleją przychody i zyski każdej firmy wraz ze wzrastającą **liczbą członków kartelu**. Stąd prawdopodobieństwo tworzenia i długotrwałej egzystencji karteli na rynkach skoncentrowanych jest większe niż na rynkach zatomizowanych. Ma to tym większe znaczenie, im większa istnieje symetria w odniesieniu do stosowanych technologii, struktury kosztów, palety produktów lub mocy produkcyjnych. Można wtedy przyjąć założenie, że przedsiębiorstwa podobne postrzegają równowagę kooperacyjną i żadne z nich nie będzie po jej osiągnięciu postrzegać swojej sytuacji jako niekorzystnej.

Posiłkując się teorią gier można wykazać, że stabilność kartelu zależy od **częstotliwości interakcji** firm. Jeżeli kartel tworzony jest na potrzeby jednorazowego wydarzenia (np. budowa stadionów piłkarskich na okoliczność olimpiady), to każde z przedsiębiorstw może pokazać się w lepszym świetle oferując

pomimo uzgodnionych cen minimalnych warunki nieco korzystniejsze, aby otrzymać dodatkowe zlecenia. Tym samym istnieje motywacja do odstąpienia od porozumień kartelowych.

Jeżeli jednak kartel ma charakter długoterminowy, ponieważ – jak w przypadku ropy naftowej – chodzi o dobra nabywane regularnie, to przedsiębiorstwo wyłamujące się z uzgodnień kartelowych musi się obawiać sankcji ze strony pozostałych członków kartelu. Zwykle jest to wykluczenie z kartelu lub krótkoterminowa wojna cenowa. W ten sposób sytuacja członka kartelu szybko staje się gorsza niż w przypadku kooperacji. Im bardziej wiarygodne są **sankcje**, tym bardziej stabilny jest kartel.

Również wysokość **czynnika dyskontującego**, czyli preferencje dotyczące przyszłości, ma znaczenie przy rozpatrywaniu ekonomicznych bodźców do tworzenia karteli. Czynniki dyskontujące zależą bezpośrednio od rynkowej stopy procentowej a pośrednio od **ryzyka**, na które narażone jest przedsiębiorstwo. Im większa rynkowa stopa procentowa, tym mniejsze znaczenie mają przyszłe przychody, co zmniejsza motywację przedsiębiorstw do koordynowania zachowań rynkowych. Podobnie jest w przypadku przedsiębiorstw narażonych są na duże ryzyko, choćby z tego powodu, że ich przyszłość uzależniona jest od sukcesu nowo wprowadzonego produktu, przez co muszą na rynku kapitałowym płacić premię za ryzyko. Równocześnie zdane są na szybkie przepływy środków pieniężnych, aby pokryć swoje zobowiązania finansowe. Także w takich sytuacjach występuje silna preferencja dla szybkich zysków i należy oczekiwać wyłamania się z porozumienia kartelowego.

Dla stabilności kartelu istotna jest ponadto **przejrzystość rynku**. Im wcześniej członkowie kartelu mogą poznać swoje ceny, koszty lub moce produkcyjne, tym mniej atrakcyjna i bardziej ryzykowna staje się opcja wyłamania się z porozumienia. Możliwość niezauważalnego wprowadzania przede wszystkim ukrytych rabatów zmniejsza się wraz z liczbą klientów, którym taki rabat będzie udzielony. Stąd rynek z dużą liczbą klientów lepiej nadaje się dla działania kartelu niż rynek wąski. Szczególnie uważnie należy badać istniejące **systemy wymiany informacji rynkowej**, zwłaszcza jawne i dorozumiane⁷¹ mechanizmy publikowania cen pod

⁷¹ Chociażby **gwarancje cenowe**, w których kupujący przyznaje nabywcy prawo do zwrotu zakupionego towaru lub zwrotu różnicy, jeżeli w określonym czasie od chwili zakupu znajdzie go w

kątem uzgodnionego postępowania uczestników rynku. Powodują one bowiem zwiększenie przejrzystości rynku, podnosząc tym samym stabilność porozumień niezgodnych z prawem konkurencji.

Znaczenie ma również rodzaj sprzedawanych dóbr. W przypadku produktów jednorodnych, można wykluczyć oprócz konkurencji poprzez ukryte rabaty także – jeszcze trudniejszą do wykrycia – konkurencję jakościową.

Na prawdopodobieństwo powstawania karteli ma również wpływ charakterystyka rynku. Do istotnych czynników zaliczyć tu można, jak przedstawiono to w rozdziale 3.4, aktualną fazę cyklu życia rynku. Na rynkach dojrzałych i schyłkowych bodźce do antykonkurencyjnej koordynacji są silniejsze niż na nowych rynkach. Dalszym czynnikiem, sprzyjającym kartelom, jest istnienie **barier dostępu do rynku**, ponieważ mimo zawyżonych cen nie trzeba się obawiać wejścia nowych konkurentów. **Elastyczność cenowa popytu** ma również znaczenia dla stabilności karteli: Tak jak w monopolu obowiązuje zasada, że zysk członka kartelu rośnie wraz z malejącą elastycznością popytu a motywacja do porozumień jest coraz silniejsza. Przy wysokiej elastyczności popytu już niewielki wzrost ceny doprowadziłby do drastycznego spadku popytu.

Można wreszcie wyciągnąć wnioski co do możliwości zaistnienia kartelu w danej branży na podstawie dostrzegalnych wzorców w kształtowaniu się wielkości ekonomicznych. Wskazówek dostarczają następujące wskaźniki, które należy zbadać w wymienionej poniżej kolejności⁷²:

Okoliczności podstawowe:

1. Niewielkie wykorzystanie mocy produkcyjnych, trwała nadwyżka popytu
2. Ścisły oligopol (np. wskaźnik koncentracji HHI > 1.000)

lepszego cenie u konkurencji. Wykorzystuje się w ten sposób własnych klientów jako źródło informacji odnośnie polityki cenowej konkurentów.

⁷² Por. Lorenz, C. (2006): Marktscreening nach Kartellstrukturen mittels ökonomischer Indikatoren: Der KMD-Kartellcheck. Opublikowane w: Wirtschaft und Wettbewerb, Nr. 7/8, str. 748-759.

Dalsze badanie; z uwagi na efektywność postępowania należy przyjąć, że dalsze badania podejmowane będą tylko w wyjątkowo istotnych przypadkach, zanim rozpocznie się poszukiwanie dowodów np. poprzez przeszukania:

3. Rzadkie i skokowe zmiany indeksu ceny nominalnej na rynku (w tym celu wartościową pomocą są urzędowe statystyki cen)
4. Wysokie stopy zwrotu na rynku
5. Dysfunkcjonalne zmiany mocy produkcyjnych
6. Niewielkie wahania udziału na rynku
7. Niewielki udział nowych produktów w obrotach
8. Niewielki wzrost produktywności

10.1.2.2 MOŻLIWOŚCI ORGANÓW OCHRONY KONKURENCJI

Po dokładnej **obserwacji rynku**, która w szczególny sposób uwzględnia charakterystykę karteli przedstawioną w powyższym rozdziale, na rynkach „zagrożonych“ organy ochrony konkurencji mogą przeprowadzać **niezapowiedziane przeszukania** przedsiębiorstw nawet przy nieokreślonym podejrzeniu o działania kartelowe i powiadamiać o tym opinię publiczną poprzez np. **intensywną współpracę z prasą**. W ten sposób dla wszystkich członków karteli na każdym badanym rynku wzrasta ryzyko wykrycia. Niepokój wewnątrz kartelu może wzrosnąć i odpowiednio restrykcyjne działania zostają podjęte, aby zapobiec odejściom choćby w ten sposób, że następują ponowne negocjacje odnośnie wielkości produkcji, mające na celu uspokojenie niezadowolonych członków kartelu. W ten sposób wzrasta prawdopodobieństwo wykrycia obciążających dowodów podczas przeszukania.

Dalszymi istotnymi zachętami dla przedsiębiorstw do wyłamywania się z karteli może być wprowadzenie **regulacji dotyczących świadka koronnego** wzgl. **programów leniency**. Zapewnia się w nich zwolnienie z karalności przedsiębiorstwa, będącego uczestnikiem porozumienia kartelowego, jeżeli pomaga ono w wykryciu kartelu. Zaufanie uczestników kartelu zostaje w ten sposób jeszcze bardziej podminowane. Zachęta do tego, aby zwracać się do organów ochrony konkurencji jest tym większa,

im wyższe są spodziewane kary i im **mniej przedsiębiorstw** w następstwie regulacji świadka koronnego skorzysta ze zwolnienia lub złagodzenia kary.

W odniesieniu do **wysokości kary** należy generalnie zwrócić uwagę na to, aby po pierwsze odebrać cały przychód dodatkowy wynikający z przynależności do kartelu (na ten temat więcej w kolejnym rozdziale). Do tego należy dodać odpowiednio wysoką kwotę zastraszającą. W kalkulacji rozsądnego członka kartelu uwzględniany jest bowiem oczekiwany skumulowany przychód dodatkowy, porównywany z iloczynem prawdopodobieństwa wykrycia kartelu i spodziewanej z tego tytułu kary. Dlatego z ekonomicznego punktu widzenia istotne jest, aby zapewnić przedsiębiorstwom „pewność co do wysokości kary“ w takim stopniu, że naliczane kary każdorazowo będą stanowiły **wielokrotność uzyskanych dodatkowych przychodów**, niezależnie od samych przychodów. Poza tym pewna „**kara minimalna**“ powinna być stosowana nawet wobec karteli działających tylko przez krótki okres. Te zasady należy opublikować a ich zastosowanie odnośnie naliczania grzywien wobec członków karteli należy **podkreślić w publikacjach**.

10.1.2.3 USTALANIE HIPOTETYCZNEJ CENY KONKURENCYJNEJ

Ustalenie hipotetycznej ceny konkurencyjnej ma decydujące znaczenie w przypadku obliczania dodatkowych przychodów wynikających z działalności kartelu, żądania wyrównania strat w przypadku naruszenia prawa konkurencji, względnie przy naliczaniu kary pieniężnej oraz w ramach postępowań w sprawie naruszeń prawa związanych z kształtowaniem cen (por. rozdział 9.3).

10.1.2.3.1 Koncepcja czasowego porównania rynku

Jedną z możliwości ustalenia wysokości dodatkowych przychodów polega na tym, że porównuje się cenę badanego dobra przed rozbiciem kartelu z ceną osiąganą po jego rozbiciu. Dane takie są łatwe do pozyskania a koncepcja jest prosta. Jej słabą stroną jest to, że zmiany struktury popytu i podaży wzgl. warunków konkurencji nie są brane pod uwagę.

Decydujące dla oceny hipotetycznej ceny konkurencyjnej w ramach koncepcji czasowego porównania rynku jest więc zbadanie, czy spadek ceny po zlikwidowaniu kartelu wynika z zakończenia funkcjonowania porozumienia kartelowego, czy też jest konsekwencją innych czynników, przede wszystkim ogólnie spadającego popytu. W celu wyjaśnienia **związku przyczynowego** należy porównać aktualną relację pomiędzy zmianami cen i popytu z analogicznym stosunkiem w okresie działania kartelu. Do przeprowadzenia takiego porównania można się posłużyć oszacowaniem elastyczności popytu podczas trwania i po zakończeniu działania kartelu oraz dowodem, że różnice pomiędzy obydwooma wyliczonymi wartościami są statystycznie istotne. Jeżeli w czasie działania kartelu ceny zmieniały się, można bez większych problemów określić elastyczność popytu w tym okresie na podstawie stwierdzonych zmian popytu. Po zakończeniu działalności kartelu dane te i tak są do ustalenia.

Należy ponadto uwzględnić, czy korekty cen w dół odpowiadają cenie konkurencyjnej, czy też wystąpiły tu nadmierne obniżki w następstwie powstania nadwyżki mocy produkcyjnych, będące konsekwencją utrzymywania się na rynku ceny kartelowej. Tak obniżone ceny nie mają długoterminowego charakteru i dlatego nie mogą być postrzegane jako ceny równowagi. Dla wyjaśnienia tej kwestii należy się przyjrzeć pozostającym do dyspozycji **mocom produkcyjnym** po zakończeniu działalności kartelu: jeżeli nie stwierdza się likwidowania mocy produkcyjnych, to zaobserwowane ceny z wystarczającym prawdopodobieństwem są cenami konkurencyjnymi. Jeżeli przeciwnie stwierdza się likwidowanie mocy produkcyjnych, należy poczekać na korektę rynku lub sięgnąć po inne metody badawcze.

10.1.2.3.2 Geograficzna i produktowa koncepcja porównawcza rynku

W tym przypadku przychód dodatkowy ustalany jest na podstawie różnicy pomiędzy ceną kartelową a ceną na rynku nieskartelizowanym, porównywalnym pod względem geograficznym lub produktowym. Decydujące znaczenie ma w tym przypadku staranna analiza **czynników dotyczących przedsiębiorstwa i struktury rynkowej**, aby znaleźć przedsiębiorstwa, których ceny są porównywalne, pomimo iż oferują swoje dobra na różnych rynkach. Albowiem w cenie odzwierciedlają się nie tylko koszty uczestniczących na rynku przedsiębiorstw, ale także egzogeniczne niedobory oraz preferencje drugiej strony rynku. W razie stwierdzenia wystarczających podobieństw pomiędzy rynkami należy ustalić ewentualne **dodatki i upusty cenowe**. Generalnie zaleca się, aby każdorazowo dla porównania służyło kilka

przedsiębiorstw. W miarę możliwości powinny to być efektywnie gospodarujące, sprywatyzowane przedsiębiorstwa.

10.1.2.3.3 Metoda oparta na kosztach

Opiera się ona na strukturze kosztów członków kartelu, która służy za podstawę do wyliczenia hipotetycznej ceny konkurencyjnej. Podstawą tej metody jest przyjęte założenie, że struktury kosztów i cen w następstwie naruszenia prawa kartelowego pozostają niezmiennie. Przedsiębiorstwa, nie poddane efektywnej konkurencji, mają niewiele bodźców do inwestowania w nowe technologie. Analogicznie nieefektywnie wygląda produkcja, która w warunkach konkurencji nie znalazłaby zbytu na rynku. Stąd **metoda** ta wydaje się być **nieprzydatna** do aproksymacji hipotetycznej ceny konkurencyjnej – nie wspominając o tym, że najpierw należałoby ustalić odpowiednie przyporządkowanie kosztów (por. r. 9.3).

10.1.2.3.4 Szacowanie cen i symulacja rynku

Hipotetyczna cena konkurencyjna ustalana jest tutaj na podstawie cen zaobserwowanych w przeszłości za pomocą **regresji wielu zmiennych** (por. r.2.2.2). Należy podkreślić, że jakość oszacowania zależy od wartości wzgl. dostępności danych. Ponadto reakcje uczestników rynku mogą być nie ujęte pomimo uwzględnienia zmian strukturalnych. Wadę tę usiłuje skorygować zastosowanie metody symulacji rynku. Opiera się ona na ekonometrycznym modelowaniu rynków oligopolistycznych przy założeniu strategicznego postępowania uczestników rynku wobec siebie. W ten sposób uwzględnia się dokładnie dynamikę strony podaży i strony popytu przy szacowaniu hipotetycznej ceny konkurencyjnej.

W istocie modele ekonometryczne kryją w sobie poważne problemy metodyczne i praktyczne, których nie można przemilczeć. Albowiem sukces zastosowania modeli ekonometrycznych uzależniony decydująco jest od tego, czy wybrany model rynku jest adekwatny, czy w ramach wybranego modelu zostaną właściwie oszacowane istotne parametry oraz czy dostępne dane spełniają wymogi ilościowe i jakościowe tak, aby na ich podstawie możliwe były dalsze wyliczenia. Analogicznie istnieją co najmniej trzy rodzaje błędów, które mogą być właściwe sugerującym matematyczną precyzję wynikom modeli. Źródła błędów mogą kumulować się i generować wynik, który nie ma wiele wspólnego z rzeczywistością (por. więcej w r. 2.2.2.2).

Główny problem polega na tym, że osoba dokonująca obliczeń nie dysponuje regularnie istotnymi danymi. Muszą one być szacowane na podstawie próbek losowych, ewentualnie nie zawsze reprezentatywnych ankiet. Zastosowanie wartości „bezpiecznych statystycznie” nie będzie zazwyczaj możliwe z powodu braku czasu i pieniędzy, natomiast stosowanie wartości szacunkowych – służących następnie jako baza kompleksowych symulacji – sądy uznałyby słusznie i zgodnie z oczekiwaniem jako podważalne, do tego prawdopodobnie znalazłyby się dalsze założenia w symulacji (specyfika rynku, parametry modelu), które można by kwestionować.

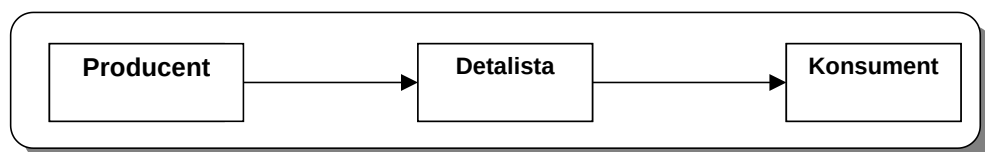
10.2 Porozumienia wertykalne

10.2.1 Założenia ogólne

W gospodarkach rynkowych zorganizowanych w oparciu o podział pracy wytwarzanie i dystrybucja dóbr, a w tym także wytwarzanie wartości dodanej, odbywa się na różnych poziomach produkcji i dystrybucji, dlatego mówi się o łańcuchu wartości dodanej pewnego produktu (względnie usługi) – poczynając od surowców poprzez półprodukty aż do produktów ostatecznych. W większości branż producenci nie sprzedają sami swoich wyrobów, lecz aby dotrzeć do klientów, korzystają z pośredniczących, hurtowych i detalicznych przedsiębiorstw handlowych. Przedsiębiorstwa na różnych poziomach łańcucha wartości dodanej często wchodzi w długoterminowe związki kontraktowe z innymi przedsiębiorstwami na rynkach znajdujących się przed (dostawcy) lub za nimi (odbiorcy) w łańcuchu produkcyjnym. Jeżeli tego typu porozumienia wertykalne zawierają zapisy, w wyniku których jedna strona kontraktu nakłada na drugą stronę, działającą na wyższym lub niższym szczeblu obrotu określone ograniczenia, mówi się o **wertykalnych ograniczeniach konkurencji**. Różne formy porozumień wertykalnych i ich następstwa można analizować przy pomocy metod ekonomicznych.

Dla unaocznienia może posłużyć przykład porozumienia wertykalnego pomiędzy producentem a detalistą, który sprzedaje jego produkty.

Wykres 25: Związek między producentem a sprzedawcą detalicznym



Producent (dostawca) i detalista (odbiorca) są niezależnymi przedsiębiorstwami, z których każde ma zamiar maksymalizować zysk. Przedsięwzięcia (np. ustalenie ceny) jednego przedsiębiorstwa niekoniecznie są równocześnie korzystne dla drugiego. Np. podwyżka ceny sprzedaży producenta oznacza większe koszty zakupu dla sprzedawcy. Z uwagi na takie relacje zwrotne każda ze stron ma ekonomiczną motywację, aby ograniczyć kontraktowo opcje postępowania drugiej strony poprzez porozumienie wertykalne. Nie zawsze tego typu porozumienie wertykalne zgodne jest z przepisami prawa konkurencji.

W europejskim prawie konkurencji przypadki porozumień wertykalnych rozstrzygane są na podstawie art. 81 ust. 1 i 3 TWE. Komisja Europejska definiuje **porozumienia wertykalne** w swoim Rozporządzeniu o Wyłączeniu Grup Wertykalnych Porozumień⁷³ jako "porozumienia lub wzajemnie uzgodnione sposoby postępowania pomiędzy dwoma lub więcej przedsiębiorstwami, z których każde prowadzi działalność celem wykonywania porozumienia na różnych poziomach produkcji lub dystrybucji, i dotyczących warunków, na których strony mogą kupować, sprzedawać lub odsprzedać określone towary lub usługi". Z reguły porozumienia wertykalne zawierane są pomiędzy nie konkurującymi ze sobą podmiotami; czasami strony kontraktu są jednak (na określonym poziomie rynku) równocześnie konkurentami (np. w przypadku franchisingodawcy, który prowadzi podwójny system sprzedaży poprzez partnerów franchisingowych i filie handlowe).

10.2.2 Niektóre istotne cechy porozumień wertykalnych

10.2.2.1 STOSUNEK PRYNCYPAŁ - AGENCI

W teorii ekonomicznej związki kontraktowe, jakie istnieją np. pomiędzy producentem a sprzedawcą detalicznym, który sprzedaje jego produkty, nazywane są często stosunkami pryncypał – agent. Oznacza to generalnie, że pryncypał zleca agentowi – z reguły w ramach długoterminowego stosunku kontraktowego – wykonywanie zadania, bez możliwości pełnej kontroli ze swojej strony. (por. także przypis dolny 24

⁷³ Rozporządzenie Komisji (WE) Nr 2790/1999 z dnia 22.12.1999 r. w sprawie stosowania art. 81 ust. 3 TWE do kategorii porozumień wertykalnych i praktyk uzgodnionych (Dz. Urz. WE L 336, z 29.12.1999), Artykuł 2 ust. 1.

i ogólnie na temat „asymetrii informacyjnej“ r. 5.3.2). Na przykład producent (pryncypał) nie może w pełni obserwować, jakie działania podejmuje sprzedawca (agent), aby sprzedać jego produkty. Teoria pryncypała–agenta (teoria agencji) analizuje – szczególnie pod kątem kompatybilności bodźców – rozwiązania problemów, które występują w specyficznym stosunku pryncypał-agent.⁷⁴

10.2.2.2 KOSZTY TRANSAKCYJNE I INTEGRACJA WERTYKALNA

Decyzje przedsiębiorstw o tym, w jaki sposób ukształtują swoje wertykalne stosunki z rynkami znajdującymi się przed lub za nimi w łańcuchu wartości dodanej, zdeterminowane są głównie kalkulacją kosztów transakcyjnych. **Koszty transakcyjne** związane są np. z koordynacją i nadzorem procesów wymiany. Zasadniczo przedsiębiorstwa (np. producenci) stają przed następującą alternatywą: albo mogą wykonać określone zadania (np. dystrybucję) w obrębie własnej hierarchii organizacyjnej (in house), albo zdać się w tym zakresie na rynek. Oznacza to, że w wielu przypadkach integracja pionowa (poprzez fuzję wertykalną (por. r. 8.2) staje się alternatywą dla porozumień wertykalnych. Zarówno w przypadku integracji wertykalnej (np. w formie sieci sprzedaży należącej do producenta) jak również w przypadku porozumień wertykalnych z innymi przedsiębiorstwami (np. franchising dystrybucyjny) generowane są – z reguły w różnej wysokości – koszty transakcyjne, które należy minimalizować. W przypadku oceny porozumień wertykalnych z punktu widzenia ochrony konkurencji powstaje kwestia, czy alternatywa połączenia wertykalnego jest dopuszczalna w świetle prawa konkurencji. Niektórzy ekonomiści domagają się, aby porozumienia wertykalne traktować niewiele surowiej niż fuzje wertykalne;⁷⁵ jest to jednak kwestia normatywna, którą musi rozwiązać ustawodawca.

10.2.2.3 KONKURENCJA MIĘDZYMARKOWA A KONKURENCJA WEWNĄTRZMARKOWA

Istotnym analitycznym rozróżnieniem z punktu widzenia analizy wpływu ograniczeń wertykalnych na konkurencję jest rozróżnienie pomiędzy konkurencją międzymarkową (**inter-brand competition**) a konkurencją wewnątrzmarkową (**intra-brand competition**). Konkurencja międzymarkowa określa konkurencję pomiędzy

⁷⁴ Patrz np. Motta (2004), str. 302 i dalsze.

⁷⁵ Patrz np. Motta (2004), str. 305.

różnymi markami (a właściwie ich producentami), natomiast konkurencja wewnątrzmarkowa jest konkurencją różnych sprzedawców w odniesieniu do produktów tej samej marki. Przy analizie określonych ograniczeń wertykalnych należy zawsze mieć na uwadze ich oddziaływanie na konkurencję międzymarkową i wewnątrzmarkową. Konkurencja wewnątrzmarkowa ograniczona jest np. poprzez sprzedaż wyłączną, konkurencja międzymarkowa osłabiana byłaby choćby poprzez np. wertykalne ograniczenie uprawnień sprzedawcy do ustalania ceny, gdyby wskutek tego ceny dla klienta ostatecznego wzrastały ponad poziom ustalony przez proces konkurencyjny, a z tego powodu zmniejszyłaby się presja konkurencyjna wywierana wzajemnie przez produktami różnych marek.

10.2.2.4 SYTUACJA EKONOMICZNA PRZEDSIĘBIORSTW

W przypadku porozumień pomiędzy konkurentami z jednej strony (porozumienia wertykalne por. r.10.1) oraz porozumień pomiędzy przedsiębiorstwami, które działają na różnych poziomach łańcucha wartości dodanej i nie konkurują ze sobą, mamy do czynienia z zasadniczo odmienną sytuacją ekonomiczną uczestniczących w nich przedsiębiorstw. Porozumienia wertykalne odnoszą się do produktów i usług, które w łańcuchu wartości dodanej są wobec siebie komplementarne (np. półprodukty i ich montaż). Już w przypadku (będącego konsekwencją wykorzystania siły rynkowej) wzrostu ceny choćby jednego półproduktu ponad poziom konkurencyjny zmniejszy się całkowity dobrobyt. W każdym razie zwiększyłaby się poprzez to (jedynie) nadwyżka producenta przedsiębiorstwa podnoszącego cenę na wyższym poziomie obrotu (upstream firm). Równocześnie zmniejszyłaby się nadwyżka przedsiębiorstwa działającego na niższym poziomie (z powodu wzrostu ceny półproduktu) przy niedoskonale elastycznym popycie. Stąd w przypadku pionowych stosunków kontraktowych interesy przedsiębiorstw przyjmują inny niż w przypadku stosunków poziomych kierunek: sprzedawca zainteresowany jest tym, aby producent ustalał niskie ceny, co zarazem wspiera ogólny dobrobyt. Z kolei producent zainteresowany jest z założenia tym, aby jego sprzedawca nie stosował zbyt wysokich cen, ponieważ przy wysokiej cenie zmniejszy się jego wielkość zbytu. Istnieje zatem obopólna zachęta do tego, aby partnerzy kontraktowi nie podnosili cen, przez co wzajemnie się dyscyplinują. Z punktu widzenia oceny konkurencji na rynku jest to pozytywne i stanowi powód do tego, aby ograniczenia wertykalne traktować często jako z założenia mniej problematyczne niż porozumienia poziome.

10.2.2.5 PODWÓJNY MONOPOLISTYCZNY NARZUT CENOWY (DOUBLE MARGINALIZATION)

Działanie samodyscyplinujące, będące cechą porozumień wertykalnych, ma jednak swoje granice i nie należy go przeceniać. Gdy jedno (lub kilka) przedsiębiorstw uczestniczących w porozumieniu wertykalnym dysponuje siłą rynkową, (por. r. 7), może wykorzystać swoją pozycję w sposób szkodliwy dla konkurencji. Można to przedstawić na przykładzie zjawiska podwójnego narzutu cenowego (ang: **double marginalization**) w ramach związku wertykalnego różnych przedsiębiorstw, przy założeniu posiadania przez nie siły rynkowej.

Zakłada się dwustopniowy związek producent – sprzedawca, jak na ilustracji 24, przyjmując również, że jedynym kosztem sprzedawcy jest cena, za jaką nabywa produkt. Oba przedsiębiorstwa, producent i sprzedawca, mają na celu maksymalizowanie swoich zysków poprzez monopolistyczny narzut cenowy na własne koszty (krańcowe). Ponieważ już cena sprzedaży hurtowej producenta zawyżona jest monopolistycznie, sprzedawca kalkuluje swoją maksymalizującą zysk cenę sprzedaży detalicznej na bazie zawyżonej ceny zakupu. W efekcie cena, jaką mają zapłacić konsumenci, w następstwie podwójnego monopolistycznego narzutu jest wyższa, niż mogłaby być, gdyby producent sam organizował (inhouse) dystrybucję swoich produktów. Wielkość sprzedaży jest mniejsza, a co za tym idzie zysk, który osiągają razem producent i sprzedawca, też jest niższy niż w przypadku producenta posiadającego własny system dystrybucji.

Pokazuje to, że łańcuch dwóch lub więcej przedsiębiorstw posiadających siłę rynkową generuje jeszcze wyższe ceny dla konsumentów ostatecznych, niż pojedyncze, zintegrowane wertykalnie przedsiębiorstwo dysponujące siłą rynkową. Przyczyną problemu double marginalization jest branie pod uwagę przez sprzedawcę (downstream firm) wyłącznie negatywnego wpływu, jaki na jego zyski ma jego własna monopolistyczna polityka cenowa, wynikającego ze spadku wielkości sprzedaży. Nie zwraca on natomiast uwagi na sprzężenie zwrotne, sprawiające, że jego polityka cenowa równocześnie ma negatywny wpływ na zysk producenta. Producent ma kilka możliwości uniknięcia wzgl. ograniczenia niekorzystnych dla niego efektów double marginalization. Jedną z możliwości byłaby wspomniana integracja wertykalna poprzez fuzję ze sprzedawcą. Jeżeli producent posiada dużą siłę rynkową także w

porównaniu ze sprzedawcą, to może próbować przejąć monopolistyczną nadwyżkę sprzedawcy poprzez odpowiednie ograniczenia wertykalne.⁷⁶

10.2.3 Ocena ograniczeń wertykalnych z punktu widzenia ich wpływu na konkurencję

Wiele ograniczeń wertykalnych nie budzi obaw w świetle prawa konkurencji i spełnia np. określone w prawie europejskim wymogi dopuszczenia zgodnie z art. 81 ust. 3 TWE. Jeżeli natomiast w jakimś pojedynczym przypadku zachodzą wątpliwości z punktu widzenia prawa konkurencji, choćby dlatego, że jakieś ograniczenie wertykalne nie spełnia wymogów rozporządzenia Komisji Europejskiej o stosowaniu art. 81 ust. 3 TWE do kategorii porozumień wertykalnych i praktyk uzgodnionych, konieczne jest osądzenie każdego przypadku indywidualnie. Dla zbadania kwestii, czy porozumienie wertykalne stanowi odczuwalne ograniczenie w rozumieniu art. 81 ust. 1 TWE, Komisja Europejska za istotne uznaje w szczególności takie czynniki, jak: pozycja rynkowa dostawcy, pozycja rynkowa konkurentów, pozycja rynkowa nabywcy, bariery dostępu do rynku, dojrzałość rynku, szczebel obrotu, właściwości produktu i inne czynniki (np. efekty kumulacji, okres obowiązywania porozumienia). Przy osądzaniu ograniczenia wertykalnego w indywidualnym przypadku należy wyważyć jego negatywne skutki z pozytywnymi.

10.2.3.1 NEGATYWNE SKUTKI OGRANICZEŃ WERTYKALNYCH

Negatywnych dla konkurencji skutków ograniczeń wertykalnych można oczekiwać szczególnie wtedy, gdy jeden lub więcej uczestników porozumienia dysponuje siłą rynkową. W takim przypadku każde przedsiębiorstwo posiadające taką siłę może przez strategiczne użycie ograniczeń wertykalnych sprawić, że konkurencja zostanie osłabiona lub wyłączona, cena będzie mogła wzrosnąć, nadwyżka konsumenta zamieni się w nadwyżkę producenta i obniży się poziom dobrobytu społecznego. Podwyżki cen mogą z kolei doprowadzić do tego, że konkurujący producenci markowi także podniosą swoje ceny, co doprowadzi także do spadku konkurencji pomiędzy markami (por. r. 10.1). Znamienne jest przy tym, że producent (pryncypał)

⁷⁶ por. Motta, jw., str. 307 i dalsze

poprzez strategiczne użycie ograniczeń pionowych deleguje wprowadzenie oczekiwanej podwyżki na swoich sprzedawców (agentów).

W szczególności można wyróżnić następujące **negatywne efekty** ograniczeń wertykalnych:

- Wykluczenie innych oferentów lub innych nabywców z rynku (**foreclosure**) szczególnie poprzez zastosowanie strategicznych barier dostępu do rynku (np. przez wyłączny system dystrybucji o odpowiednim działaniu);
- Ograniczenie konkurencji pomiędzy markami (np. wertykalna polityka cenowa może sprzyjać zmawianiu się oferentów marek);
- Ograniczenie konkurencji wewnątrzmarkowej (np. przez selektywną dystrybucję).

Niniejsze negatywne, w konsekwencji zmniejszające dobrobyt skutki mogą być uwarunkowane przez różne ograniczenia wertykalne. Równocześnie ograniczenia wertykalne mogą także powodować wymienione poniżej skutki pozytywne.

10.2.3.2 POZYTYWNE SKUTKI OGRANICZEŃ WERTYKALNYCH

W określonych przypadkach należy ocenić pozytywnie z ekonomicznego punktu widzenia skutki ograniczeń wertykalnych, co też znajduje uzasadnienie w prawie konkurencji. Ograniczenia wertykalne mogą przyczynić się do usunięcia problemów koordynacji rynkowej i uczynić rynki bardziej efektywnymi:

- Często wymienianym problemem w dystrybucji jest występowanie pozytywnych efektów zewnętrznych (por. r. 5.1) wskutek działań dystrybucyjnych jednego sprzedawcy, z których korzystają (za darmo) także inni sprzedawcy (problem gapowicza). Jeżeli inni sprzedawcy korzystają jako gapowicze np. z dobrego doradztwa handlowego określonego sprzedawcy drogich urządzeń elektronicznych, w konsekwencji prowadzi to zbyt niskiego poziomu usług dodatkowych. Rozwiązaniem problemu gapowicza pomiędzy sprzedawcami może być na przykład wyłączność sprzedaży.
- Czasami istotne przy wprowadzaniu nowych produktów jest to, aby były one sprzedawane wyłącznie przez określonych „wysokiej jakości sprzedawców“, którzy przy nieograniczonej dystrybucji nie oferowaliby nowych produktów, ponieważ inni sprzedawcy mogliby za darmo korzystać z wizerunku **wysokiej**

jakości sprzedawcy (problem gapowicza). W tym kontekście (przejściowo) system wyłącznej lub selektywnej dystrybucji może być uzasadniony.

- Aby w dystrybucji **umożliwić** niezbędne specyficzne **inwestycje początkowe**, w przypadku wchodzenia na nowe, regionalne rynki, wyłączność na dany obszar może być uzasadniona.
- Wertykalne ograniczenia mogą **wywoływać** także **problemy lock in (pochwycenia)** (por. r. 5.3.2). Przykładowo, dostawca półproduktów musi zainwestować w specyficzne urządzenia produkcyjne lub szkolenia dla swoich odbiorców. Jeżeli nie będą mu przyznane szczególne warunki kontraktu, inwestycja nie zostanie wykonana lub zostanie wykonana w zbyt małym zakresie.
- **Problem lock in** występuje również w przypadku **przenoszenia** specyficznego **know how**, które raz przekazane nie może już zostać odebrane z powrotem. Może to uzasadniać np. zakazy konkurencji wewnątrz systemu franchisingu.
- **Zabezpieczenie jednolitości i jakości** asortymentu u sprzedawców w ramach systemu franchisingu może uzasadniać zobowiązanie do zakupów. Gdyby franchisingodawca nie mógł zapewnić jednolitości i jakości asortymentu, byłoby w takim przypadku zagrożone dobre imię systemu franchisingu, a w konsekwencji także jego szanse rynkowe.
- W celu **osiągnięcia korzyści skali** (por. r. 5.2.2) w **dystrybucji** może być czasem uzasadnione ograniczenie przez producenta liczby sprzedawców firmowych (np. poprzez selektywny system dystrybucji).
- Inne przypadki, np. jeżeli producenci dają swoim sprzedawcom firmowym kapitał (pożyczkę), aby **wyrównać niedoskonałości rynków kapitałowych**, ograniczenia wertykalne jako zabezpieczenia są uzasadnione.

10.2.4 Poszczególne formy ograniczeń wertykalnych

Ograniczenia wertykalne występują w bardzo różnorodnych formach, pojedynczo i w pakietach. Poniżej przedstawiono przykłady selektywnej dystrybucji, wertykalnych cen wiążących, franchising i umowa wyłączności.

Selektywne umowy dystrybucyjne (**selektywna dystrybucja**) ograniczają, podobnie jak w przypadku wyłącznej dystrybucji, liczbę autoryzowanych punktów sprzedaży

oraz możliwości dalszej odsprzedaży. Ograniczenie liczby sprzedawców uzależnione jest od określonych kryteriów doboru, które wiążą się z rodzajem sprzedawanego produktu. Wyróżnia się **jakościową** i **ilościową** dystrybucję selektywną.⁷⁷ W konsekwencji ograniczona zostaje sprzedaż nieautoryzowanym sprzedawcom, a w ten sposób zmniejsza się konkurencja wewnątrzmarkowa. Systemy selektywnej dystrybucji mogą być tylko wtedy uzasadnione jako działanie podnoszące efektywność, gdy właściwości produktu wymagają selektywnej dystrybucji. Ilościową dystrybucję selektywną należy ocenić pozytywnie, gdy (dodatkowo) umożliwia korzyści skali w dystrybucji. Jakościowa dystrybucja selektywna może np. uzasadniać spadek konkurencji wewnątrzmarkowej na tyle, na ile jest niezbędna, aby zapewnić określony poziom jakości i doradztwa.

Jako **wertykalne zobowiązania cenowe** określa się porozumienia lub uzgodnione sposoby postępowania, które nakazują odbiorcy jakiegoś produktu (pośrednio lub bezpośrednio) stosowanie stałej lub minimalnej ceny przy odsprzedaży. Wertykalne porozumienia cenowe ograniczają zarówno konkurencję wewnątrzmarkową jak i konkurencję pomiędzy markami. Z punktu widzenia producenta produktów markowych występuje bodziec do wertykalnego zobowiązania cenowego, jeżeli pojawi się problem podwójnego narzutu cenowego. Artykuł 4 lit. a) Rozporządzenia w sprawie wyłączenia porozumień wertykalnych i praktyk uzgodnionych KE definiuje **ograniczenie uprawnień kupującego do ustalania własnej ceny sprzedaży** jako restrykcję nieobjętą wyłączeniem. Jako mniej problematyczne dla konkurencji zaklasyfikowano **zalecane ceny sprzedaży** – o ile nie są równoważne minimalnej lub stałej cenie sprzedaży. Ceny zalecane mogą jednak mieć skutki dławiące konkurencję (np. z powodu ich funkcji informacyjnej dla sprzedawców) oraz sprzyjać zmonopolizowaniu dostawców produktu.

Franchising zawiera system uzgodnień kontraktowych (umowy franchisingu), za pomocą których franchisingodawca (FD) przenosi know how na niezależne przedsiębiorstwa – franchisingobiorca (FB), a biorcy franchisingu sprzedają produkty lub usługi pod marką franchisingodawcy, przy czym stosują jego model biznesowy i

⁷⁷ Cechą odróżniającą jest zastosowanie czysto jakościowych kryteriów doboru (np. wymogi serwisu, przeszkolenie sprzedawców) z jednej strony, wobec ilościowych kryteriów doboru (np. ograniczenie liczby punktów sprzedaży, minimalne obroty) w drugim przypadku.

korzystają z akcji reklamowych.⁷⁸ Relacje franchisingowe obejmują szereg ograniczeń wertykalnych. Ograniczenia te znajdują uzasadnienie, o ile konieczne są dla zabezpieczenia przekazania know how przez dawcę franchisingu oraz dla zapewnienia jednolitości i jakości towarów wzgl. Asortymentu, itp.⁷⁹ Szczególnie, kiedy dawca franchisingu (w jakimś stopniu) dysponuje siłą rynkową, niektóre elementy umów franchisingu mogą budzić zastrzeżenia w świetle konkurencji (np. art. 81 ust. 1 TWE lub art. 82 TWE).

Umowy na wyłączność to np. **zobowiązania do zakupu**.⁸⁰ Zobowiązują one kupującego do skoncentrowania jego popytu w uzgodnionym okresie w znacznej części lub w całości u jednego dostawcy, co może skutkować zamknięciem rynku. Tego typu zobowiązania z reguły uzasadnione są tylko w przypadku drogich, nieodwracalnych inwestycji, ponieważ w takich przypadkach tworzą się wzajemne zależności (efekt **lock in**). Jeżeli w takich okolicznościach brakuje możliwości sankcji wynikających z kontraktu, które zmuszałyby strony kontraktu do lojalnego postępowania, nie powstawałyby użyteczne, specyficzne inwestycje wzgl. nie dochodziłoby do transakcji lub rozmiar ich byłby zbyt mały (por. r. 5.3.2). Z punktu widzenia prawa konkurencji zobowiązania do zakupu w zasadzie tylko wtedy budzą zastrzeżenia, gdy dostawca dysponuje siłą rynkową. Z tego powodu zobowiązania do zakupu często omawia się w kontekście nadużycia pozycji dominującej na rynku w ramach dyskusji o unilateralnych ograniczeniach konkurencji.

⁷⁸ Dobrą wiedzę wprowadzającą w temat franchisingu daje wydany przez Metzlaiff, Carsten (Hrsg.), *Praxishandbuch Franchising*, München 2003.

⁷⁹ Podstawowe kryteria oceny systemów franchisingowych według prawa europejskiego zostały ustalone przez ETS w decyzji w sprawie Pronuptia z dnia 28.01.1986 r.

⁸⁰ Również wyłączność dostaw jest umową wyłączności. Umowa o wyłączności dostaw zobowiązuje dostawcę do dostarczania produktu wyłącznie jednemu określönemu odbiorcy. Poprzez to inni odbiorcy mogą zostać wyłączeni z rynku.

Literatura

van Bergeijk, P. (ed.) (2005): Modelling European mergers: theory, competition policy and case. Cheltenham pp.

Kokkoris, Ioannis:

Do merger simulation and critical loss analysis differ under the SLC and dominance test? w: European competition law review; 27(2006) Heft 5 ; str. 249-260.

Lorenz, Christian:

Marktscreening nach Kartellstrukturen mittels ökonomischer Indikatoren : der KMD-Kartellcheck; Hinweise auf das Vorliegen kartellierter Strukturen in Form typischer Muster von Marktprozessen über sogenannte collusive marker. w: Wirtschaft und Wettbewerb ; 56(2006) 7/8 ; str. 748-759.

International Competition Network (2005): ICN Investigative Techniques Handbook for Merger Review; dostępne pod adresem:

http://www.internationalcompetitionnetwork.org/media/library/conference_4th_bonn_2005/Investigative_Techniques_Handbook.pdf .

International Competition Network (2005): ICN Merger Guidelines Workbook; dostępne pod adresem:

http://www.internationalcompetitionnetwork.org/media/library/conference_5th_capetown_2006/ICNMergerGuidelinesWorkbook.pdf .

Maddala, G.S. i Lajal Lahiri (2001): Introduction to Econometrics, Wiley & Sons, 3. Aufl., Chichester.

Dougherty, Christopher (2006): Introduction to Econometrics, Oxford University Press, Oxford.

Binmore, Ken (1991): Fun and Games. A Text on Game Theory, Houghton Mifflin Company, Lexington.

Chamberlain, E. H. (1933) : The Theory of Monopolistic Competition.

Myerson, Roger B. (1997): Game Theory: Analysis of Conflict, Harvard University Press, Cambridge.

Osborne, Martin J. (2003): An Introduction to Game Theory, Oxford University Press, Oxford.

Gibbons, Robert (1992): A Primer in Game Theory, Financial Times Print Int., ##.

Epstein, Roy J. i Daniel L. Rubinfeld (2004): Merger Simulation with Brand-Level Margin Data: Extending PCAIDS with Nests, Advances in Economic Analysis & Policy, Vol. 4, Bd. 2.

Baker, Jonathan B. i Daniel L. Rubinfeld (1999): Empirical Methods in Antitrust Litigation: Review and Critique, American Law and Economics Review, Vol. 5, str. 386 - 435.

Nalebuff, Barry (2004): Bundling as a Way to Leverage Monopoly, Yale School of Management Working Paper, dostępne pod adresem:

<http://ssrn.com/abstract=586648>.

Greenlee, Patrick et al (2004): An Antitrust Analysis of Bundled Loyalty Discounts, US DoJ Antitrust Division Economic Analysis Group, dostępne pod adresem:

<http://ssrn.com/abstract=600799>.

Evans i Padilla (2004): Excessive Pricing: Using Economics to Define Administrable Legal Rules, CEMFI Working Paper No. 0416.

Motta, Massimo (2004): Competition Policy, Cambridge University Press, Cambridge.

POZYCJA DOMINUJĄCA

1. Harold Demsetz (1974), 'Two Systems of Belief About Monopoly'
2. Franklin M. Fisher i John J. McGowan (1983), 'On the Misuse of Accounting Rates of Return to Infer Monopoly Profits'

KONCENTRACJE

3. Oliver E. Williamson (1968), 'Economies as an Antitrust Defense: The Welfare Tradeoffs'
4. Kenneth G. Elzinga i Thomas F. Hogarty (1973), 'The Problem of Geographic Market Delineation in Antimerger Suits'
5. Robert D. Willig (1991), 'Merger Analysis, Industrial Organization Theory, and Merger Guidelines'

6. Barry C. Harris i Joseph J. Simons (1989), 'Focusing Market Definition: How Much Substitution is Necessary?'

7. Jonathan B. Baker i Timothy F. Bresnahan (1988), 'Estimating the Residual Demand Curve Facing a Single Firm'

8. Gregory J. Werden i Luke M. Froeb (1994), 'The Effects of Mergers in Differentiated Products Industries: Logit Demand and Merger Policy'

POROZUMIENIA HORYZONTALNE

9. George J. Stigler (1964), 'A Theory of Oligopoly'

10. George A. Hay i Daniel Kelley (1974), 'An Empirical Survey of Price Fixing Conspiracies'

11. Thomas E. Cooper (1986), 'Most-Favored-Customer Pricing and Tacit Collusion'

OGRANICZENIA WERTYKALNE

12. Lester G. Telser (1960), 'Why Should Manufacturers Want Fair Trade?'

13. Benjamin Klein i Kevin M. Murphy (1988), 'Vertical Restraints as Contract Enforcement Mechanisms'

DRAPIEŻNICTWO CENOWE

14. John S. McGee (1958), 'Predatory Price Cutting: The Standard Oil (N.J.) Case'

15. Robert H. Bork (1978), 'Injury to Competition: The Law's Basic Theories'

16. Paul Milgrom i John Roberts (1982), 'Predation, Reputation, and Entry Deterrence'

OGRANICZANIE DOSTĘPU DO RYNKU

17. Steven C. Salop i David T. Scheffman (1987), 'Cost-Raising Strategies'

18. Timothy J. Brennan (1988), 'Understanding "Raising Rivals' Costs"'

19. Janusz A. Ordover, Garth Saloner i Steven C. Salop (1990), 'Equilibrium Vertical Foreclosure'

20. David Reiffen (1992), 'Equilibrium Vertical Foreclosure: Comment'

21. Philippe Aghion i Patrick Bolton (1987), 'Contracts as a Barrier to Entry'
22. David A. Butz i Andrew N. Kleit (2001), 'Are Vertical Restraints Pro- or Anticompetitive? Lessons from Interstate Circuit'

SIECIOWE EFEKTY ZEWNĘTRZNE

23. Michael L. Katz i Carl Shapiro (1985), 'Network Externalities, Competition, and Compatibility'
24. Paul A. David (1985), 'Clio and the Economics of QWERTY'
25. S.J. Liebowitz i Stephen E. Margolis (1990), 'The Fable of the Keys'

www.uokik.gov.pl
