

BANK I KREDYT

Czasopismo NBP poświęcone ekonomii i finansom
National Bank of Poland's Journal on Economics and Finance

maj
2008

- 3 **Piotr Popowski, Agnieszka Sawicka**
Wybrane aspekty oddziaływania niepewności na decyzje inwestycyjne polskich przedsiębiorstw. Wyniki badania empirycznego
 Some Aspects of the Impact of Uncertainty on Investment Decisions Made by Polish Enterprises: Results of an Empirical Study
- 21 **Piotr Bańbuła**
Options and Market Expectations: Implied Probability Density Functions on the Polish Foreign Exchange Market
 Opcje a oczekiwania rynku: estymacja i wykorzystanie implikowanych funkcji gęstości prawdopodobieństwa na polskim rynku walutowym
- 36 **Aleksandra Małek**
Obligacja jako narzędzie sekurytyzacji ryzyka śmiertelności
 Bond as a Tool of Mortality Risk Securitization
- 50 **Agnieszka Cenkier**
Partnerstwo publiczno-prywatne – dlaczego warto zabiegać o jego rozwój
 Public-Private Partnership – Why Should We Insist on its Development?
- 62 **Adam Koronowski**
Peter L. Bernstein, *Capital Ideas Evolving*
 Review of the book by Peter L. Bernstein, *Capital Ideas Evolving*
- 65 **Tomasz Tokarski**
Andrzej Rzońca, *Czy Keynes się pomylił? Skutki redukcji deficytu w Europie Środkowej*
 Review of the book by Andrzej Rzońca, *Was Keynes Wrong? Consequences of Deficit Reduction in the Central Europe*

Europejska Integracja Monetarna (educational insert in Polish only)

Leokadia Oręziak

Polityka budżetowa w strefie euro

Budget Policy in the Euro Area

Rada Naukowa/Scientific Council

Peter Backé (Oesterreichische Nationalbank), Wojciech Charemza (University of Leicester), Stanisław Gomułka (London School of Economics and Political Science), Marek Góra (Szkoła Główna Handlowa), Marek Gruszczyński (Szkoła Główna Handlowa), Urszula Grzełowska (Szkoła Główna Handlowa), Danuta Hübner (European Commission), Krzysztof Jajuga (Akademia Ekonomiczna we Wrocławiu), Bartłomiej Kamiński (University of Maryland; The World Bank), Jerzy Konieczny (Wilfrid Laurier University), Wojciech Maciejewski (Uniwersytet Warszawski), Krzysztof Marczewski (Szkoła Główna Handlowa; Instytut Badań Rynku, Konsumpcji i Koniunktury), Ewa Miklaszewska (Uniwersytet Ekonomiczny w Krakowie), Timothy P. Opiela (DePaul University, Chicago), Witold Orłowski (Niezależny Ośrodek Badań Ekonomicznych; Szkoła Biznesu Politechniki Warszawskiej), Zbigniew Polański (zastępca przewodniczącego/Deputy Chairman, Narodowy Bank Polski; Szkoła Główna Handlowa), Bogusław Pietrzak (Szkoła Główna Handlowa; Narodowy Bank Polski), Wiesława Przybylska-Kapuścińska (Akademia Ekonomiczna w Poznaniu), Zbyněk Revenda (Vysoká škola ekonomická v Praze), Michel A. Robe (American University; U.S. Commodity Futures Trading Commission), Michał Rutkowski (The World Bank), Sławomir Stanisław Skrzypek (przewodniczący/Chairman, prezes/President, Narodowy Bank Polski), Adalbert Winkler (European Central Bank), Charles Wyplosz (Graduate Institute of International Studies, Geneva)

Kolegium Redakcyjne/Editorial Board

Piotr Boguszewski, Tomasz Chmielewski, Elżbieta Czarny, Krzysztof Gajewski (sekretarz kolegium redakcyjnego/Assistant Editor), Małgorzata Iwanicz-Drozdowska, Ryszard Kokoszczynski, Adam Koronowski, Wojciech Pacho, Bogusław Pietrzak (zastępca redaktora naczelnego/Deputy Managing Editor), Zbigniew Polański (redaktor naczelny/Managing Editor), Andrzej Rzońca, Cezary Wójcik, Zbigniew Żółkiewski

Zgodnie z wykazem sporządzonym przez Ministerstwo Nauki i Szkolnictwa Wyższego dla potrzeb przyszłej oceny parametrycznej jednostek naukowych, publikacjom naukowym w „Banku i Kredycie” przyznawane jest 6 punktów.

Wersje elektroniczne artykułów publikowanych w „Banku i Kredycie” są dostępne za pośrednictwem serwisu *Social Science Research Network* (<http://www.ssrn.com>)

Electronic versions of the articles published in "Bank i Kredyt" are available at the *Social Science Research Network* (<http://www.ssrn.com>)

Wydawca/Publisher

Narodowy Bank Polski

Kontakt/Contact

ulica Świętokrzyska 11/21, 00-919 Warszawa, Poland

tel.: +48 22 585 43 26

fax: +48 22 826 99 35

e-mail: krzysztof.gajewski@mail.nbp.pl

<http://www.nbp.pl/bankikredyt>

Projekt/Project

DOCTORAD

Skład i Druk/Typesetting and printing

Drukarnia NBP/Printing House of the NBP

Korekta/Editing

Departament Komunikacji Społecznej NBP/Department of Information and Public Relations NBP

Prenumerata/Subscription

„RUCH” SA - wpłaty na prenumeratę przyjmują: jednostki kolportażowe właściwe dla miejsca zamieszkania lub siedziby prenumeratora (dostawa w sposób uzgodniony). Wpłaty przyjmuje Oddział Krajowej Dystrybucji Prasy „RUCH” SA na konto: Pekao SA IV O/Warszawa 12401053-40060347-2700-401112-001 lub kasa Oddziału. Cena prenumeraty ze zleceniem dostawy za granicę jest o 100% wyższa od krajowej. Zlecenia na prenumeratę dewizową, przyjmowane od osób zamieszkających za granicą, realizowane są od dowolnego numeru w danym roku kalendarzowym. Wpłaty są przyjmowane na okresy kwartalne w terminie: do 5.12 - na I kw. następnego roku, do 5.03 - na II kw.br., do 5.06 na III kw. br., do 5.09 na IV kw. br. Informacje o warunkach prenumeraty w „RUCH” SA OKDP, ul. Jana Kazimierza 31/33 00-958 Warszawa, można uzyskać pod tel. 532-87-31, 532-88-20.

www.ruch.pol.pl

Prenumerata własna i zamawianie pojedynczych egzemplarzy:

Narodowy Bank Polski - Departament Komunikacji Społecznej,

ulica Świętokrzyska 11/21, 00-919 Warszawa,

nakład: 1200

konto:

Centrala NBP - Departament Operacyjno-Rachunkowy

nr konta: NBP DOR 871010-0000-0000-1323-9600-0000

2008 r. - 204,00 zł, 1 egz. - 17,00 zł

Wybrane aspekty oddziaływania niepewności na decyzje inwestycyjne polskich przedsiębiorstw. Wyniki badania empirycznego

Some Aspects of the Impact of Uncertainty on Investment Decisions Made by Polish Enterprises: Results of an Empirical Study

*Piotr Popowski, Agnieszka Sawicka**

pierwsza wersja: 20 grudnia 2007 r., ostateczna wersja: 7 maja 2008 r., akceptacja: 12 maja 2008 r.

Streszczenie

W niniejszym artykule zbadano wpływ nieodwracalności inwestycji oraz siły konkurencji na zależność pomiędzy inwestycjami a niepewnością. W analizie empirycznej wykorzystano dane jednostkowe pochodzące z badania ankietowego obejmującego ponad 800 polskich przedsiębiorstw niefinansowych, przeprowadzonego przez Narodowy Bank Polski w 2006 r.

Wykazano, że stopień odwracalności inwestycji wpływa na zależność pomiędzy niepewnością a inwestycjami – gdy inwestycje są nieodwracalne, wpływ niepewności na inwestycje staje się negatywny. Wpływ pozycji konkurencyjnej przedsiębiorstwa na siłę oddziaływania niepewności na inwestycje okazał się w badanej próbie statystycznie nieistotny.

Słowa kluczowe: inwestycje, niepewność, nieodwracalność

Abstract

The paper investigates the impact of the irreversibility of investments and competition on the sign of the relationship between investment and uncertainty. The empirical analysis uses firm-level data and is based on a survey of over 800 Polish non-financial companies carried out by the National Bank of Poland in 2006.

We demonstrate that the relationship between investment and uncertainty is influenced by the extent to which investments are irreversible. The results indicate that in the case of irreversibility the relationship between uncertainty and investment becomes negative. However, the degree of market power the investor enjoys in the product market turned out to be insignificant for the investigated relationship.

Keywords: firm investment, uncertainty, irreversibility

JEL: D81, D92, C24

* Narodowy Bank Polski, Instytut Ekonomiczny, e-mail: piotr.popowski@mail.nbp.pl, agnieszka.sawicka@mail.nbp.pl. Poglądy i opinie wyrażone w tekście są poglądami i opiniami autorów i niekoniecznie są zbieżne ze stanowiskiem NBP.

1. Wstęp

Wpływ niepewności na decyzje inwestycyjne przedsiębiorstw często jest przedmiotem badań w literaturze. Z punktu widzenia teorii nie ulega wątpliwości, że występuje związek między niepewnością a inwestycjami podmiotów gospodarczych. Jednak wyniki dotychczasowych badań empirycznych – szczególnie przeprowadzanych na poziomie danych jednostkowych – nie są w pełni zgodne nie tylko co do siły, ale i co do kierunku oddziaływania. Jednym ze spójnych wniosków z tych prac jest natomiast to, że wyniki badań relacji między niepewnością a inwestycjami silnie zależą od przyjętych założeń, w tym szczególnie od sposobu pomiaru – złożonej, lecz nieobserwowalnej z natury – niepewności. Rezultat ten podkreśla znaczenie badań koniunktury, umożliwiających uzyskanie subiektywnych ocen niepewności, formułowanych bezpośrednio przez inwestorów. Wykorzystanie takich ocen pozwala zatem w pewnym stopniu na uniknięcie, częstego w podobnych pracach, problemu przyjmowania zbyt restrykcyjnych (w rzeczywistości rzadko spełnionych) założeń umożliwiających kwantyfikację niepewności.

Celem niniejszych badań jest empiryczne określenie kierunku wpływu niepewności na decyzje inwestycyjne polskich przedsiębiorstw. Spośród wielu obszarów, w których może oddziaływać niepewność, w niniejszym badaniu skoncentrowano się na niepewności dotyczącej popytu na produkty przedsiębiorstwa, uznawanej w literaturze za jedno z najistotniejszych źródeł niepewności. Na podstawie wyników dotychczasowych prac można się spodziewać, że **wzrost niepewności może hamować inwestycje w przedsiębiorstwach**. Ze względu na potwierdzone w literaturze podstawowe znaczenie założeń i pewnych uwarunkowań dla uzyskiwanych wyników w niniejszym badaniu uwzględniono m.in. poziom cenowej elastyczności funkcji popytu (na produkty przedsiębiorstwa), wyrażonej przez siłę konkurencji na rynku produktów danego podmiotu, oraz stopień **odwracalności inwestycji**, czyli skalę trudności z odsprzedaniem zainstalowanego majątku. Zgodnie z teorią konkurencja ma wpływ na kierunek relacji między niepewnością a inwestycjami. **W przypadku przedsiębiorstw działających w warunkach doskonałej konkurencji niepewność nie hamuje inwestycji, lecz może je stymulować**. Jest to tzw. **efekt Hartmana-Abła**. Hartman (1972) oraz Abel (1983) twierdzili, że przy pewnych założeniach¹ w przedsiębiorstwach działających w warunkach konkurencji wzrost niepewności może zwiększać inwestycje. Większa niepewność (np. wariacji cen) powoduje bowiem wyższy oczekiwany zysk z krajowej jednostki kapitału, co zachęca przedsiębiorców do inwestowania (Carruth et al. 1998, s. 2). W odwrotnej sytuacji, tzn. w przypadku zaburzeń mechanizmów rynko-

wych, należy oczekiwać, że wzrost niepewności będzie obniżał inwestycje (Fuss, Vermuelen 2004, s. 2). Analogiczny wpływ na relacje pomiędzy niepewnością a inwestycjami może mieć charakter inwestycji. Gdy przedsiębiorstwo inwestuje w majątek, który jest stosunkowo łatwo zbywalny na rynku wtórnym (nie ma asymetrii kosztów dostosowań), wyższa niepewność nie powinna negatywnie oddziaływać na inwestycje. Dotychczasowe badania pokazują ponadto, że **wpływ niepewności na inwestycje jest negatywny, jeśli realizowane inwestycje cechują się wysokim stopniem nieodwracalności**.

Powodem podjęcia niniejszych badań było m.in. to, że badania na poziomie danych jednostkowych, wykorzystujące subiektywne podstawy oceny niepewności są w literaturze nieliczne. Według wiedzy autorów istnieją tylko dwie podobne prace, w których zastosowano analogiczną metodę pomiaru niepewności. Poniżej prezentujemy pierwsze takie badanie dla polskich przedsiębiorstw. Dzięki możliwości wykorzystania unikatowych danych jednostkowych, którą dają badania realizowane przez NBP, niepewność oszacowano na podstawie rozkładu subiektywnych prognoz popytu badanych przedsiębiorstw. Do budowy oceny niepewności wykorzystano dane z pytania ankietowego, w którym każde z przedsiębiorstw oceniało, jakie są szanse, że dynamika przyszłej sprzedaży przyjmie wartość z kolejnych, zdefiniowanych przedziałów. W ten sposób powstały indywidualne rozkłady prawdopodobieństw, a odchylenia standardowe tych rozkładów wykorzystano jako miarę niepewności.

Rozpoznanie wpływu niepewności na decyzje przedsiębiorstw może być istotne dla działania mechanizmów transmisji pieniężnej, jest więc ważne dla skuteczności polityki gospodarczej. Niepewność może bowiem tłumaczyć słabsze oddziaływanie instrumentów tej polityki na sferę realną (Bloom et al. 2006).

Przedstawione w dalszej części badania opiera się na jednostkowych danych przedsiębiorstw, uzyskanych w ramach rocznego badania ankietowego zrealizowanego w NBP w 2006 r.

Narodowy Bank Polski rozpoczął swoje badania ankietowe przedsiębiorstw w 1995 r. w celu rozpoznania i monitorowania sytuacji ekonomicznej i koniunktury w sektorze przedsiębiorstw, w tym szczególnie tych jej elementów, które mają podstawowe znaczenie dla prowadzenia polityki pieniężnej banku centralnego. Badanie to wraz z realizowanym równoległe badaniem kwartalnym składa się na system badań przedsiębiorstw w NBP i jest jednym z ważniejszych źródeł informacji zarówno o bieżącej, jak i prognozowanej sytuacji przedsiębiorstw. Uczestnictwo w tych badaniach jest dobrowolne i nieodpłatne. Badanie realizowane jest za pośrednictwem 16 oddziałów okręgowych NBP, które kontaktują się (najczęściej drogą korespondencyjną, ale także bezpośrednio) z wytypowanymi do badań przedsiębiorstwami i prowadzą wywiad według sporządzonego

¹ Założenia te to m.in. neutralne nastawienie inwestorów do ryzyka, funkcja produkcji o stałych korzyściach skali.

w Centrali NBP formularza ankietowego. Próba obejmuje obecnie około 800 przedsiębiorstw niefinansowych ze wszystkich działów PKD i sektorów własności.

W wyniku zaprezentowanych poniżej analiz wykazano, że **odwracalność inwestycji zmniejsza negatywne oddziaływanie niepewności na inwestycje**. Odwracalność inwestycji ma także decydujący wpływ na znak zależności pomiędzy niepewnością a inwestycjami. Wykazano bowiem, że niepewność tłumi inwestycje w przypadku przedsiębiorstw inwestujących w majątek trudno zbywalny (tzn. gdy występuje silna asymetria kosztów dostosowań), natomiast jeśli inwestycje są odwracalne, niepewność jest neutralna dla inwestycji.

Teza mówiąca, że występowanie silnej konkurencji na rynku produktów przedsiębiorstwa zmniejsza negatywne oddziaływanie niepewności na inwestycje, nie znalazła potwierdzenia w niniejszym badaniu.

2. Oddziaływanie niepewności na decyzje inwestycyjne przedsiębiorstw – przegląd teorii

2.1. Nurty badań

Wpływ niepewności na decyzje inwestycyjne przedsiębiorstw jest przedmiotem badań od lat 60. ubiegłego wieku².

Pierwszy z dwóch głównych nurtów badań zapoczątkowały prace Hartmana (1972) i Abła (1983). Analizowano w nich decyzje inwestycyjne przedsiębiorstw o neutralnym nastawieniu do ryzyka (krańcowa produktywność kapitału jest wówczas wypukłą funkcją zmiennych, których kształtowanie jest obciążone niepewnością) i działających w warunkach doskonałej konkurencji. Przyjęto także założenie, że funkcję produkcji cechują niemalejące korzyści skali i nie występuje problem nieodwracalności inwestycji. Przy takich uwarunkowaniach **większa niepewność zwiększa krańcową zyskowność kapitału i pobudza inwestycje**.

Drugi nurt badań wiąże się z **teorią opcji** i narodził się dzięki odpowiedniemu zastosowaniu tej teorii do analizy decyzji inwestycyjnych przedsiębiorstw. Punktem wyjścia było tu założenie, że decyzje inwestycyjne przedsiębiorstwa można wycenić tak jak decyzje w zakresie inwestycji finansowych. Do rozwoju tego nurtu przyczynili się m.in. McDonald i Siegel (1986), Dixit i Pindyck (1994) oraz Abel i Eberly (1994). W badaniach uwzględniono nieodwracalność inwestycji. Nieodwracalność inwestycji można rozpatrywać jako formę kosztów dostosowań, a wynika ona z faktu, że pewna część kosztów zakupu nowych składników kapitału ma charakter kosztów utopionych. Koszty te są asymetryczne,

ponieważ redukcja kapitału nie pozwala na odzyskanie nakładów poniesionych na jego zainstalowanie. Z nieodwracalnymi inwestycjami w majątek mamy do czynienia wówczas, gdy poziom asymetrii jest wysoki i inwestor nie ma możliwości odsprzedaży tego majątku lub może go odsprzedać jedynie znacznie poniżej ceny zakupu. Wykazano, że jeśli w warunkach niepewności inwestycje są nieodwracalne, a przedsiębiorstwa mogą decydować o momencie realizacji inwestycji (tzw. *timing*), wówczas przedsiębiorstwo „zachowuje opcję” realizacji inwestycji w przyszłości i wstrzymuje inwestycje bieżące. Wskutek odroczenia planów inwestycyjnych przedsiębiorstwo ponosi ryzyko utraty przyszłych zysków, ale zyskuje czas potrzebny na zdobycie dodatkowych informacji, które mogą przyczynić się do zmniejszenia poziomu niepewności i umożliwić podjęcie korzystniejszej dla przedsiębiorstwa decyzji. **Wraz ze wzrostem niepewności wartość opcji się zwiększa, co skłania przedsiębiorstwa do wstrzymywania się z inwestycjami, a to z kolei prowadzi do obniżenia bieżących inwestycji**. Dodatnia wartość opcji powoduje, że pojawia się rozbieżność pomiędzy klasyczną wartością bieżącą projektu inwestycyjnego (NPV) a uwzględniającą niepewność kalkulacją bieżącej wartości projektu inwestycyjnego dla inwestora. Oznacza to, że aby inwestycja mogła być zrealizowana, jej NPV musi być istotnie większa od zera, żeby pokryć straty wynikające z opóźnienia inwestycji i utrzymania opcji jej realizacji w przyszłości. Ważnym skutkiem tego rozumowania jest istnienie **wartości progowej** (*threshold effect*). Jeśli stopa zwrotu z inwestycji przewyższa tę wartość, przedsiębiorstwo podejmuje inwestycje, a jeśli jest poniżej wartości progowej – wstrzymuje się z inwestycjami. Niepewność zwiększa dystans między krańcową produktywnością kapitału uzasadniającą podjęcie inwestycji a krańcową produktywnością kapitału uzasadniającą dezinvestycje. W ten sposób zwiększa się przedział, w którym inwestycje są zerowe, ponieważ przedsiębiorstwa wolą stosować strategię „poczekamy, zobaczymy” niż podejmować kosztowne inwestycje o nieprzewidywalnych konsekwencjach (Bloom et al. 2006). Nieodwracalność inwestycji ma więc wpływ na kształtowanie się inwestycji w przedsiębiorstwach i powoduje, że przynajmniej na poziomie poszczególnych projektów inwestycje nie mają przebiegu ciągłego, lecz raczej skokowy, z częstszymi okresami zerowej aktywności inwestycyjnej (Butzen et al. 2002, s. 4).

Ta cecha procesu inwestycyjnego jest niezwykle istotna pod względem makroekonomicznym, gdyż może prowadzić do pewnej sztywności procesu akumulacji kapitału (Carruth et al. 1998, s. 2) i histerezy inwestycji. To zjawisko może z kolei wyjaśniać mniejszą skuteczność instrumentów polityki gospodarczej prowadzonej w warunkach niepewności (stymulowanie inwestycji może wówczas dawać gorsze wyniki).

² Zauważono wówczas, że niedoskonała konkurencja decyduje o kierunku zmian majątku przedsiębiorstwa na niepewność co do zmian popytu. Prace w latach 70. dotyczyły m.in. wpływu niepewności na poziom produkcji w firmach charakteryzujących się awersją do ryzyka (Guiso, Parigi 1999, s. 188).

Warto podkreślić, że modele tworzone w ramach teorii opcji z nieodwracalnymi inwestycjami w warunkach niepewności nie dają odpowiedzi na temat optymalnego poziomu inwestycji w przedsiębiorstwie. Umożliwiają natomiast identyfikację tych czynników, które mogą wpływać na wartość progową, krytyczną dla decyzji o realizacji inwestycji (Carruth et al. 1998, s. 8).

2.2. Problemy związane z badaniem wpływu niepewności na inwestycje

Znaczenie założeń

Z teoretycznego punktu widzenia nie ma wątpliwości, że występuje **istotny związek między kształtowaniem się niepewności a inwestycjami przedsiębiorstw**. Wyniki dotychczasowych badań empirycznych **nie są jednak w pełni zgodne** nie tylko co do **siły**, ale i **znaku tej relacji** (Fuss, Vermuelen 2004, s. 1). Rozbieżności pojawiają się szczególnie pomiędzy wynikami badań realizowanych na poziomie danych jednostkowych; bardziej spójne są wnioski z analiz danych zagregowanych.

Przyczyną uzyskiwania niejednoznacznych wyników, jeśli chodzi o **kierunek** związku pomiędzy niepewnością a inwestycjami, jest **wrażliwość tego związku na przyjmowane założenia**. Wzrost niepewności może stymulować inwestycje lub ograniczać je w zależności od przyjętej kombinacji założeń co do technologii produkcji (czyli efektów skali, *returns to scale*), konkurencyjności na rynku produktów przedsiębiorstwa (czyli kształtu funkcji popytu), stopnia odwracalności inwestycji (czyli właściwości krzywej kosztów dostosowań, *adjustment cost*, czy skłonności menadżerów przedsiębiorstwa do ryzyka (Fuss, Vermuelen 2004, s. 2; Guiso, Parigi 1999, s. 185).

W przypadku przedsiębiorstwa o dużej sile monopolistycznej, dysponującego technologią charakteryzującą się malejącymi korzyściami skali, którego majątek jest stosunkowo trudno zbywalny na rynku wtórnym, bardziej prawdopodobne jest, że wyższa niepewność będzie tłumić inwestycje. W sytuacji odwrotnej, tj. gdy podmiot działa w warunkach silnej konkurencji, inwestuje w majątek, który można relatywnie łatwo upłynnić i dysponuje technologią o rosnących korzyściach skali, należy oczekiwać dodatniego wpływu niepewności na inwestycje.

W literaturze często zakłada się, że przedsiębiorcy są neutralnie nastawieni do ryzyka. Jak pokazują analizy teoretyczne przeprowadzone przez Bernanke (1983), Caballero (1991) oraz Dixita i Pindycka (1994), nawet przy założeniu braku awersji do ryzyka kierunek i siła wpływu niepewności na inwestycje mogą być różne w zależności od pozostałych czynników: odwracalności inwestycji, siły konkurencji oraz korzyści skali. Również

w niniejszym artykule, podobnie jak w większości analogicznych badań, nie uwzględniono nastawienia do ryzyka jako trudnej do zmierzenia indywidualnej cechy każdego przedsiębiorstwa, lecz raczej założono brak awersji do ryzyka.

Jak badać nieodwracalność inwestycji?

Guiso i Parigi (1999) zaproponowali dwie możliwości ujęcia odwracalności w modelu empirycznym.

Według pierwszej metody przyjmuje się, że większą odwracalnością charakteryzują się inwestycje przedsiębiorstw, które leasingowały elementy majątku trwałego albo kupowały lub sprzedawały na rynku wtórnym używane elementy majątku trwałego. Podstawową wadą tego podejścia jest to, że nie pozwala ono na rozróżnienie firm, których majątek jest relatywnie trudno zbywalny, i tych, które nie miały potrzeby leasingować i kupować lub sprzedawać na rynku wtórnym elementów majątku trwałego.

Druga metoda, pozbawiona tej wady, polega na badaniu korelacji pomiędzy sytuacją badanej firmy i całego sektora, do którego ona należy. Podstawą tego podejścia jest określenie stopnia odwracalności inwestycji, a więc płynności majątku, jako różnicy pomiędzy wartością jego elementów (przy założeniu możliwości ich wykorzystania w procesie wytwórczym) a ceną oferowaną przy ewentualnej odsprzedaży. Gdy firma ma nadmierne moce produkcyjne lub jest w złej sytuacji finansowej i zdecyduje się na odsprzedaż elementów majątku wytwórczego, to prawdopodobnie znajdzie kupca wśród firm o podobnym profilu działalności. Jeśli jednak impuls, który skłonił daną firmę do sprzedaży majątku, dotknął także potencjalnych kupców, będą oni skłonni zaoferować niższą cenę. Guiso i Parigi badali skorelowanie wielkości produkcji badanej firmy z wielkością sprzedaży w odpowiedniej sekcji. Jeśli jest ono silne, to znaczy, że dominują impulsy wspólne dla całej branży i majątek danej firmy jest trudno zbywalny. W przeciwnym wypadku, gdy wahania sprzedaży są typowe dla danej jednostki i nieskorelowane z odpowiednią wielkością dla reszty podmiotów, majątek danej firmy należy uznać za względnie łatwo zbywalny.

W celu przybliżenia odwracalności w niniejszej pracy inwestycji wykorzystano informację o tym, jak przedsiębiorstwa subiektywnie oceniają możliwości odsprzedaży elementów majątku trwałego. Do badanych przedsiębiorstw skierowano bezpośrednie pytanie na temat możliwości zbycia składników majątku firmy (ramka 1).

Zakres nieodwracalności inwestycji w polskich przedsiębiorstwach

Wyniki badania NBP pokazują, że **nieodwracalność inwestycji**, rozumiana jako trudności z odsprzedażą

majątku produkcyjnego (maszyn i urządzeń) wykorzystywanego do wytwarzania głównego produktu³, **często występuje w badanej próbie przedsiębiorstw** – łącznie dotyczyła około **połowy respondentów** (około – ponieważ część firm nie potrafiła ocenić sytuacji na rynku wtórnym, por. wykres 1).

W tej grupie około 18% inwestorów oceniło majątek jako niezbywalny m.in. ze względu na jego silną specjalizację. W przypadku 34% firm majątek jest trudno zbywalny, gdyż ewentualna odsprzedaż wiązałaby się albo z koniecznością zaakceptowania niższej ceny, albo z długotrwałym poszukiwaniem kupca. Jedynie około 14% przedsiębiorstw oceniło, że majątek produkcyjny firmy jest łatwo zbywalny, tzn. można relatywnie szybko znaleźć kupca oferującego satysfakcjonującą cenę zakupu. Pozostała, bardzo liczna grupa firm (około 34% firm) nie miała wiedzy na temat sytuacji na rynku wtórnym dóbr kapitałowych.

Zbadano odwracalność inwestycji w zależności od wielkości przedsiębiorstwa (aby wykluczyć wpływ profilu działalności firmy na łatwość zbycia majątku, badanie to przeprowadzono w grupie przedsiębiorstw prze-

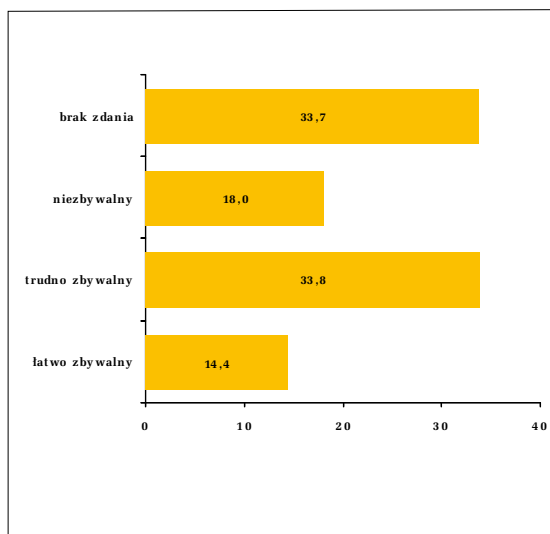
twórstwa przemysłowego⁴). Wyniki pokazują, że **mniejsze przedsiębiorstwa zdecydowanie słabiej orientują się w możliwościach odsprzedaży swojego majątku**, co również może wpływać na bardziej ostrożne planowanie inwestycji (wykres 5). Łącznie 30% firm MSP (zatrudniających do 249 pracowników), w tym ponad połowa firm najmniejszych (do 50 pracowników), nie potrafiło ocenić możliwości odsprzedaży własnego majątku produkcyjnego (19% w grupie dużych firm, por. wykres 5). W badanej próbie małe firmy rzadziej realizują inwestycje (w 2006 r. znacznych inwestycji dokonało 58% firm z sektora MSP i 74% dużych podmiotów), co mogłoby tłumaczyć ich słabszą wiedzę na temat możliwości upłynnienia majątku na wtórnym rynku. Jest to jednak tylko częściowe wytłumaczenie, ponieważ pewna różnica między stanem wiedzy w sektorze MSP i dużych firmach na temat możliwości odsprzedaży posiadanego majątku utrzymuje się także po wyłączeniu podmiotów nierealizujących inwestycji.

Wśród firm mających rozeznanie w możliwościach sprzedaży majątku można zaobserwować – co jest zgodne z intuicją – że **majątek dużych firm jest nieco trudniej zbywalny niż firm z sektora MSP**. Trudności z odsprzedażą majątku bądź brak możliwości jego zbytu przewiduje 85% firm dużych i 76% firm z sektora MSP. Relacje te kształtują się podobnie w zależności od pozycji rynkowej (wykres 6). Ogólnie **wraz z pozycją rynkową firm rośnie ich świadomość możliwości odsprzedaży majątku, a majątek w większym stopniu jest nieodwracalny**.

³ Kategoria majątku przedsiębiorstwa jest pojemna i zależy m.in. od specyfiki działalności danego przedsiębiorstwa. W niniejszym badaniu ocena nieodwracalności dotyczy majątku produkcyjnego wykorzystywanego do produkcji głównego produktu przedsiębiorstwa. Dzięki temu uzyskano przynajmniej częściową porównywalność tej kategorii w poszczególnych przedsiębiorstwach. Wśród czynników, które wiążą się z problematyką nieodwracalności inwestycji w majątek przedsiębiorstw, trzeba wymienić wiek posiadanego majątku oraz stopień jego zużycia. Ze względu na ograniczoną pojemność formularza ankietowego zagadnienia te nie zostały ujęte w niniejszym badaniu. W przyszłości podobne badania powinny jednak zostać rozszerzone o wpływ tych czynników na zakres odwracalności inwestycji przedsiębiorstw

⁴ Ogólne wnioski nie zmieniają się także w pełnej próbie przedsiębiorstw, obejmującej wszystkie działy PKD.

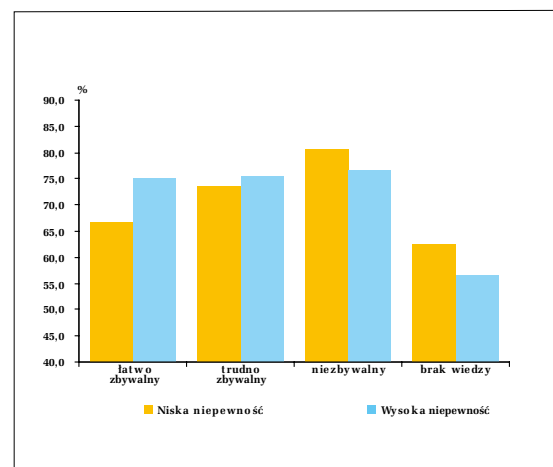
Wykres 1. Ocena odwracalności inwestycji*



* Możliwość odsprzedaży na rynku wtórnym majątku wykorzystywanego do produkcji głównego produktu przedsiębiorstwa

Źródło: badanie ankietowe przedsiębiorstw, NBP.

Wykres 2. Udział przedsiębiorstw planujących duże inwestycje w klasach wydzielonych ze względu na stopień odwracalności inwestycji i poziom niepewności*



* Wysoka niepewność – współczynnik różnicowania prognoz popytu powyżej średniej dla próby, niska – poniżej średniej.

Źródło: badanie ankietowe przedsiębiorstw, NBP.

Niepewność może skłaniać do wstrzymywania nieodwracalnych inwestycji, gdyż przedsiębiorstwa mogą przedkładać utrzymywanie niedostatecznych zdolności produkcyjnych nad ponoszenie ryzyka utworzenia i utrzymywania nadwyżkowego kapitału. Badania NBP przeprowadzone na danych przekrojowych pokazały, że przedsiębiorstwa **niemające możliwości odsprzedaży majątku są mniej skłonne do podejmowania inwestycji, jeśli działają w warunkach relatywnie wyższej niepewności** (wykres 2).

W przypadku przedsiębiorstw mających łatwy dostęp do rynków wtórnych wyższa niepewność zwiększa aktywność inwestycyjną (może więc nawet zachęcać do podejmowania inwestycji). Podobne wyniki – pokazujące, że wzrost niepewności może wpływać na zwiększenie inwestycji, jeśli są one odwracalne – dotyczą belgijskich przedsiębiorstw i zostały zaprezentowane w pracy (Cassimon et al. 2002, s. 18). Obserwacja ta potwierdza zatem tezę, że w warunkach niepewności przedsiębiorstwa wolą ograniczać inwestycje, jeśli mają one charakter nieodwracalny. Jest to spójne z wnioskami płynącymi z modelu zaprezentowanego w rozdziale 3.

Znaczenie pozycji konkurencyjnej na rynku produktów

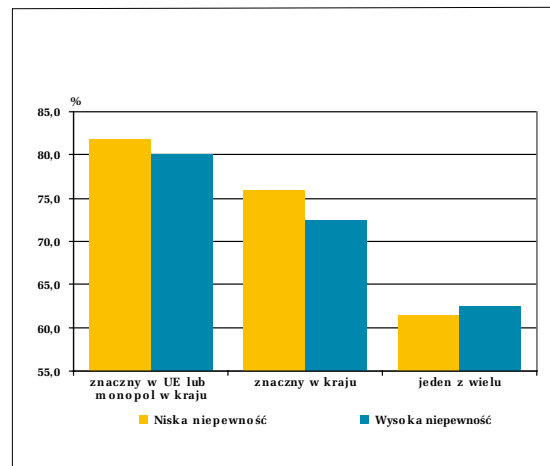
Teoria decyzji inwestycyjnych w warunkach niepewności mówi, że dla danych technologii i przy asymetrycznej funkcji kosztów dostosowań wpływ niepewności zależy od stopnia konkurencyjności na rynku produktów. Im większa jest siła monopolowa przedsiębiorstw, tym bardziej prawdopodobne, że wzrost niepewności zmniejszy inwestycje.

Do oceny pozycji konkurencyjnej przedsiębiorstw, która jest przybliżeniem cenowej elastyczności funkcji popytu, wykorzystano teorię mówiącą, że efektywność danego podmiotu jest dodatnio skorelowana z jego pozycją rynkową (Lerner 1934; Guiso, Parigi 1999). Pozycję rynkową danej firmy można więc określić na podstawie oceny jej rentowności na tle pozostałych przedsiębiorstw. Guiso i Parigi uznali, że za firmy z silną pozycją rynkową należy uznać te, których rentowność jest wyższa od średniej dla danej sekcji.

Zmienną wykorzystaną w niniejszym badaniu, określającą pozycję rynkową przedsiębiorstwa, zbudowano na podstawie danych ankietowych. Przedsiębiorstwa same oceniały swoją siłę rynkową: czy są monopolistami, mają silną pozycję na rynku, czy też są jedną z wielu firm w danej branży i ich pozycja rynkowa jest słaba. Jedną z przesłanek wybrania subiektywnej podstawy wyznaczenia pozycji rynkowej był fakt, że pozwala ona na uniknięcie problemów związanych z doborem zmiennych do określenia tej pozycji, np. zmiennych ilościowych.

Wyniki badania ankietowego NBP nie potwierdzają tezy, że niepewność może mieć negatywny wpływ na aktywność inwestycyjną w warunkach niedoskonałej

Wykres 3. Udział przedsiębiorstw planujących duże inwestycje w klasach wg pozycji rynkowej w zależności od ocen poziomu niepewności*



*Wysoka niepewność – współczynnik zróżnicowania prognoz popytu powyżej średniej dla próby, niska – poniżej średniej.

Źródło: badanie ankietowe przedsiębiorstw, NBP.

konkurencji. Wprawdzie dane przekrojowe pokazują, że **wyższa niepewność zmniejsza aktywność inwestycyjną w grupie przedsiębiorstw, które były istotnymi producentami lub monopolistami** (wykres 3). Jednak różnice aktywności inwestycyjnej w badanych przekrojach nie są istotne w sensie statystycznym. Spójne z tymi wnioskami są także wyniki modelu zaprezentowane w rozdziale 3.

Problem pomiaru niepewności

Wyniki badań nad wpływem niepewności na inwestycje są wrażliwe nie tylko na przyjęte założenia, ale także na **sposób pomiaru – złożonej i nieobserwowalnej z natury – niepewności**. Niepewność jest jednym z elementów składających się na warunki gospodarowania przedsiębiorstw. W rzeczywistości przedsiębiorstwa doświadczają niepewności co do wielu czynników: m.in. wielkości popytu, poziomu cen, płac, wysokości podatków, regulacji prawnych, poziomu stóp procentowych i kursu walutowego, przebiegu postępu technologicznego.

W literaturze stosuje się **wiele różnych metod szacowania niepewności** i wśród ekonomistów **nie ma konsensusu**, która z nich jest najlepsza.

Można wyróżnić kilka sposobów ujmowania niepewności w modelach empirycznych. Klasyfikacji miar niepewności można dokonać ze względu na dwie cechy: **przedmiot, którego niepewność dotyczy**, oraz **charakter wykorzystanych danych i metodę analitycznej reprezentacji** owej zmiennej.

Jeśli chodzi o pierwszy podział – **podział przedmiotowy** – należy wyodrębnić niepewność w kształ-

towaniu się wielkości popytu, cen produktów przedsiębiorstwa lub cen czynników produkcji. Tak skonstruowane miary mają silne umocowanie w modelach teoretycznych, jednak żadna z nich nie ujmuje całej niepewności, z jaką mają do czynienia firmy. Alternatywą może być modelowanie przychodów ze sprzedaży lub wyników finansowych badanych podmiotów. Zaletą tego rozwiązania jest łączne ujęcie niepewności co do zmian czynników stojących zarówno po stronie popytu, jak i produkcji (cen produktów i wielkości popytu, jak również zmian cen czynników wytwórczych, zmian technologicznych czy gustów konsumentów). Może jednak być krytykowane za brak odpowiedniej podstawy teoretycznej (Ghosal, Loungani 2000).

Spśród metod szacowania niepewności można wymienić rozwiązanie zaproponowane przez Pindycka (1986). Jego podstawą jest wariancja dziennych stóp zwrotu z akcji przedsiębiorstw. Metoda ta opiera się na założeniu, że na zmianę stóp zwrotu wpływają między innymi oczekiwane przez rynek zmiany popytu i cen czynników produkcji. Taka miara może jednak nie do końca odzwierciedlać niepewność, jaką odczuwają osoby podejmujące decyzje o inwestycjach. Ceny akcji nie kształtują się bowiem wyłącznie pod wpływem czynników fundamentalnych, lecz dodatkowo podlegają wahaniom na skutek np. nieracjonalnych zachowań inwestorów, przepływów kapitału spekulacyjnego i występowania bąbli spekulacyjnych.

Istotnym problemem związanym z tworzeniem empirycznych miar niepewności jest zatem **dobór zmiennych**. Większość miar niepewności jest liczona na podstawie zmiennych obserwowalnych – danych pochodzących ze sprawozdań finansowych. Takie podejście wymaga założenia, że wszystkie analizowane przedsiębiorstwa sporządzają swoje prognozy cen, popytu, wyników itd. (które są podstawą do wyznaczenia niepewności), używając tego samego algorytmu, co jest założeniem zbyt mocnym. Indywidualne zbiory informacji dostępnych dla menedżerów poszczególnych firm, wykorzystywane do planowania i podejmowania decyzji o przeprowadzeniu inwestycji, są bowiem różne i często daleko wykraczają poza ten uwzględniony w ogólnym modelu ekonometrycznym. Dlatego należy się spodziewać, że tak skonstruowana miara niepewności będzie przeszacowana.

W większości badań niepewność szacowana jest **pośrednio jako pochodna określonej kombinacji zmiennych**. Jest jednak wiele zastrzeżeń wobec takiego podejścia. Po pierwsze występuje tu problem estymacji i specyfikacji modelu. Po drugie korelacja inwestycji z niepewnością opartą na miarach zmienności, zwłaszcza na poziomie zagregowanym, może być także wynikiem korelacji z innymi czynnikami fundamentalnymi, pominiętymi w modelu. Istnieje też wątpliwość, czy miara niepewności powinna być *backward-* czy *forward-looking*. Za koncepcją *forward-looking* przemawia to, że

pozwała ona ująć więcej informacji o zakresie i poziomie niepewności postrzeganej przez inwestorów (Carruth et al. 1998, s. 23).

W drugiej grupie metod szacowania niepewności można wskazać metody, w których wykorzystuje się **miary niepewności oparte na wynikach badań jakościowych i subiektywnych ocenach przedsiębiorstw**. To podejście pozwala na uniknięcie konieczności przyjmowania założeń odnośnie do zakresu i zawartości zbioru informacji, którymi dysponują osoby podejmujące decyzje o inwestycjach.

Rodzajem takiej miary jest **współczynnik Theila**, który konstruuje się na podstawie oczekiwań przedsiębiorstw co do kierunku zmian cen lub popytu na produkty przedsiębiorstwa. Zaletą tego podejścia jest stosunkowo duża dostępność danych potrzebnych do skonstruowania miary niepewności. Do tego celu wykorzystuje się odsetki przedsiębiorstw deklarujących odpowiedni wariant odpowiedzi na pytania z testów koniunktury. Przedsiębiorstwa oceniają sygnały płynące z gospodarki i na ich podstawie podają najbardziej prawdopodobny scenariusz kształtowania się popytu (wzrost, spadek lub brak zmiany). Na podstawie „wariancji” tych odpowiedzi przybliża się wariancję rozkładu prawdopodobieństwa wystąpienia szokowych zmian popytu dla całej gospodarki. Główną wadą tego wskaźnika jest agregacja danych i wynikająca z niego konieczność przyjęcia założenia o homogeniczności niepewności w poszczególnych grupach przedsiębiorstw. Większość modeli teoretycznych zakłada natomiast, że decyzje **inwestycyjne przedsiębiorstw są bardziej wrażliwe na szoki specyficzne, oddziałujące na poziomie poszczególnych przedsiębiorstw, niż na szoki o szerszym zakresie oddziaływania**. Jest to zatem jedna z ważniejszych przesłanek wyboru metod szacowania niepewności na poziomie jednostkowym, również w przypadku niniejszego badania.

Kolejna miara w klasie ocen subiektywnych jest wyznaczona na podstawie rozkładu **indywidualnych prawdopodobieństw spełnienia się prognoz dotyczących określonych kategorii ekonomicznych**. W literaturze prognozy takie dotyczą np. kształtowania się popytu na wyroby przedsiębiorstwa lub cen tych wyrobów. Zaletą takiego podejścia jest to, że ocena taka uwzględnia podejście *forward-looking*, co jest uzasadnione z teoretycznego punktu widzenia. Decyzje inwestycyjne dotyczą bowiem nieznannej przyszłości i uwzględniają oczekiwania przedsiębiorstwa. Bieżące inwestycje są zatem odzwierciedleniem sytuacji przedsiębiorstwa w przeszłości i jego ówczesnych oczekiwań. Miary oparte wyłącznie na danych finansowych nie zawierają elementu *forward-looking*. Podstawową przewagą tej miary nad miarą wykorzystującą współczynnik Theila jest ponadto, że do-

starcza ona informacji o niepewności na poziomie poszczególnych podmiotów, a nie tylko branży. Poważnym ograniczeniem rozwoju tego nurtu badań są jednak **mała dostępność danych** i wysokie koszty pozyskania danych jednostkowych. W praktyce powoduje to, że badania wykorzystujące taką zmienną ograniczają się do analizy przekrojowej.

Autorom niniejszego artykułu znane są jedynie dwa badania (Guiso, Parigi 1999, Patillo 1998), w których wykorzystano bezpośrednio miary niepewności formułowane przez przedsiębiorstwa. W obu badaniach miary niepewności skonstruowane zostały na podstawie uzyskanych w ankiecie subiektywnych rozkładów prawdopodobieństw oczekiwanego popytu.

Dzięki wyjątkowej możliwości wykorzystania danych jednostkowych, jaką oferują badania NBP, analogiczną metodę szacowania niepewności można było zastosować także w niniejszym badaniu.

Oczywiste jest, że przedsiębiorstwa podejmują decyzje o zwiększeniu lub zmniejszeniu albo modernizacji potencjału wytwórczego, opierając się na oczekiwaniach co do kształtowania się czynników zasadniczych dla prowadzonej przez nie działalności. Nie wymaga to założenia, że przedsiębiorstwa muszą wiedzieć, jakie są rozkłady prawdopodobieństw realizacji poszczególnych scenariuszy kształtowania się tych czynników (w tym przypadku wielkości sprzedaży). Należy jednak uznać, że na podstawie dostępnych informacji badane podmioty są w stanie oszacować owe prawdopodobieństwa, a także uwzględniają te szacunki (subiektywne rozkłady prawdopodobieństw), planując inwestycje.

W odpowiedzi na pytanie ankietowe przedsiębiorstwa podawały, jakie jest prawdopodobieństwo, że stopa wzrostu sprzedaży (w 2006 r.) znajdzie się w każdym z przedziałów wymienionych w pytaniu. W ten sposób uzyskano informację o subiektywnych rozkładach prawdopodobieństw zmian wielkości sprzedaży. Na podstawie owych rozkładów dla każdego podmiotu policzono średnią⁵ oraz odchylenie standardowe oczekiwanej stopy wzrostu sprzedaży. Miary dyspersji indywidualnego rozkładu prawdopodobieństw zmian popytu zostały użyte do oszacowania poziomu niepewności co do kształtowania się wielkości sprzedaży na poziomie przedsiębiorstwa (część formularza ankietowego zawierającego to pytanie zaprezentowano w aneksie, ramka 1).

Jedną ze **słabości metody** szacowania niepewności zastosowanej w niniejszym badaniu może być fakt, że uwzględnia ona **tylko jedno ze źródeł niepewności**, tzn. niepewność dotyczącą popytu. Argumentem za wyborem popytowego ujęcia niepewności jest fakt, że znaczenie popytowego źródła niepewności jest często podkreślane w literaturze.

Rozkłady indywidualnych miar niepewności w obszarze popytu oraz średnich oczekiwanych wielkości zmian popytu zamieszczono w ramce 2 w aneksie.

Problem pomiaru inwestycji

Kolejnym problemem w badaniach nad wpływem niepewności na inwestycje jest **sposób ujęcia zmiennej objaśnianej – czyli wydatków inwestycyjnych firm**. W większości prac wykorzystuje się w tym celu bieżące wydatki inwestycyjne, które są publikowane w sprawozdaniach finansowych przedsiębiorstw. Tymczasem z ekonomicznego punktu widzenia bardziej odpowiednią kategorią są inwestycje planowane. Decyzje inwestycyjne podejmowane są bowiem przez przedsiębiorstwa z pewnym wyprzedzeniem, wyłącznie na podstawie dostępnych wówczas, i obciążonych niepewnością, informacji. Bieżące inwestycje są zaś skutkiem decyzji podejmowanych na podstawie informacji dostępnych w przeszłości. Dla właściwego pomiaru niepewności wskazane by było zatem, żeby wszystkie informacje, które firma uwzględnia w tym procesie i które są włączone do modelu, także były mierzone jak najbliżej momentu podejmowania decyzji inwestycyjnej w przedsiębiorstwie (Butzen et al. 2002, s. 1). Takie podejście, w którym wykorzystuje się inwestycje planowane, zastosowano w pracach: Pattillo (1998); Guiso, Parigi (1999); Butzen (2002); Ninh et al. (2003) oraz w niniejszym opracowaniu.

W prezentowanym badaniu zmienną określającą planowane inwestycje przedsiębiorstw skonstruowano na podstawie odpowiedzi na pytanie o zamierzenia inwestycyjne przedsiębiorstw na najbliższy rok oraz skalę tych przedsięwzięć (pytanie zaprezentowano w aneksie, ramka 1). Informacja ta została następnie przekształcona w zmienną binarną określającą, czy dany przedsiębiorca planował w 2006 r. realizację istotniejszych projektów inwestycyjnych. Charakter zmiennej objaśnianej wymaga zastosowania modelu probabilistycznego, którego interpretacja pozwala na stwierdzenie, jaki wpływ na prawdopodobieństwo planowania inwestycji przez przedsiębiorstwo mają poszczególne zmienne objaśniające. Taka konstrukcja zmiennej objaśnianej nie pozwala na określenie ilościowego wpływu niepewności na wielkość inwestycji, jednak umożliwia wnioskowanie o istnieniu zależności i jej ewentualnym kierunku, co wystarczy do zweryfikowania hipotez postawionych w niniejszym opracowaniu.

Wpływ ograniczenia finansowego

Realizując duże projekty inwestycyjne, przedsiębiorstwa najczęściej korzystają z finansowania zewnętrznego. Jeżeli nie ma możliwości przeprowadzenia inwestycji przy wykorzystaniu jedynie środków własnych, uzyskanie finansowania zewnętrznego jest decydujące dla realizacji inwestycji.

⁵ Średnie wykorzystano w równaniu inwestycji w celu wyizolowania wpływu, jaki na decyzje inwestycyjne wywiera oczekiwana wielkość zmiany poziomu sprzedaży.

Ponieważ instytucje udzielające finansowania przedsiębiorstwom mają mniej informacji do oceny projektu inwestycyjnego i kondycji przedsiębiorstwa ubiegającego się o finansowanie zewnętrzne (występuje tzw. problem asymetrii informacji), może się okazać, że bank przeszacował ryzyko związane z inwestycją, co nie pozwala na udzielenie kredytu. W rezultacie przedsiębiorstwo nie pozyska środków niezbędnych do przeprowadzenia inwestycji, która z jego punktu widzenia byłaby opłacalna. W takim przypadku inwestycja nie jest realizowana z powodu ograniczenia finansowego (*financial constraints*). Ujęcie tego czynnika w równaniu inwestycji pozwoli na uwzględnienie skutku ograniczenia finansowego wynikającego z asymetrii informacji.

Ograniczenie finansowe jest często określane na podstawie wpływu bieżącej kondycji finansowej (reprezentowanej przez relację *cash-flow* do kapitału) na inwestycje przedsiębiorstwa, tzn. istotny dodatni wpływ relacji *cash-flow* do kapitału na inwestycje świadczy o ograniczeniu finansowym przedsiębiorstwa (Fuss 2004; Ghosal, Loungani 2000; Guiso, Parigi 1999). Owa miara może jednak również reprezentować różnice efektywności działania, a nie tylko wpływ czynników finansowych na inwestycje; dlatego niektórzy badacze zdecydowali się na użycie miar alternatywnych.

Ghosal i Loungani (2000) podzielili firmy na małe i duże. Mimo że sam rozmiar przedsiębiorstwa nie determinuje dostępności finansowania zewnętrznego, to jest silnie skorelowany z innymi czynnikami, których wpływ jest istotny. Wysoki poziom asymetrii informacji, który podnosi koszty finansowania zewnętrznego, dotyczy bowiem głównie podmiotów młodych, a więc niemających historii kredytowej, niedysponujących majątkiem pod zabezpieczenie kredytów i podlegających wahaniom niewynikającym z ogólnej sytuacji w danej branży, czyli najczęściej stosunkowo małych firm.

Guiso i Parigi zbudowali na podstawie danych ankietowych zmienną binarną, która pozwala na określenie, czy firma miała ograniczony dostęp do finansowania zewnętrznego (czy był racjonowany kredyt). Dostępność odpowiednich informacji umożliwiła wykorzystanie podobnej zmiennej w niniejszym analizie.

Problem agregacji danych

W badaniach nad wpływem niepewności na inwestycje dominują prace prowadzone na poziomie **danych zagregowanych**, czemu sprzyja **większa dostępność** tego typu danych. Agregacja danych ma jednak poważne konsekwencje w postaci **ryzyka neutralizowania się efektów** występujących na poziomie jednostkowym. Sytuacja ta może wystąpić zwłaszcza wówczas, gdy różnego rodzaju zaburzenia kształtujące niepewność w badanych grupach przebiegają wielokierunkowo. W literaturze wymienia się dwie możliwe przyczyny tego, że zagregowany efekt raczej się znosi, niż uśrednia (Bernanke

1983, s. 85–106). Po pierwsze przedsiębiorstwa, podejmując różnego typu decyzje, uwzględniają globalnie oddziałujące czynniki makroekonomiczne, takie jak niepewność co do poziomu stóp procentowych, kursów walut, inflacji, szoków monetarnych i fiskalnych, czy zmiany regulacji prawnych. Po drugie zagregowana niepewność może być wywołana przez indywidualne jednostki, a więc na poziomie mikroekonomicznym. Jeśli bowiem w przedsiębiorstwie utrzymuje się niepewność co do trwałości zagregowanych szoków (czy szok jest trwały, czy przejściowy) lub zasięgu ich oddziaływania (czy szok zagregowany będzie odczuwany na poziomie jednostkowym), przedsiębiorstwo może opóźniać swoje decyzje inwestycyjne do czasu napływu nowych informacji obniżających tę niepewność. Ze względu na nieodwracalność inwestycji początkowo przejściowe szoki mogą się jednak przekształcać w szoki permanentne (Carruth et al. 1998, s. 8; Fuss et al. 2004, s. 11).

Przegląd dotychczasowych badań pokazuje, że **badania na danych zagregowanych dają bardziej spójne wyniki niż badania na danych jednostkowych** (Carruth et al. 1998, s. 14). Pomimo znacznego różnicowania stosowanych metod empirycznych, typów wykorzystywanych modeli oraz sposobów szacowania niepewności wyniki badań realizowanych na poziomie danych zagregowanych ogólnie wskazują na **negatywne oddziaływanie niepewności na inwestycje**.

Większość argumentów (poza gorszą dostępnością danych jednostkowych) przemawia jednak za tym, że badania wpływu niepewności na inwestycje **powinny być prowadzone na poziomie jednostkowym**. Po pierwsze dlatego, że decyzje inwestycyjne podejmowane są na poziomie przedsiębiorstwa. Po drugie, analiza danych jednostkowych umożliwia **uwzględnienie czynników specyficznych** dla przedsiębiorstwa, które – zgodnie z literaturą dotyczącą nieodwracalności – **mają większe znaczenie dla wyjaśniania zachowań inwestycyjnych niż czynniki globalne**, oddziałujące na wszystkie przedsiębiorstwa (Carruth et al. 1998, s. 15).

W niniejszym badaniu wykorzystano unikatowe w skali polskiej gospodarki dane jednostkowe, uzyskane dzięki szczegółowym badaniom sektora przedsiębiorstw prowadzonym w NBP.

3. Wpływ niepewności na decyzje inwestycyjne przedsiębiorstw – wyniki badania empirycznego dla polskich przedsiębiorstw

W powyższej części opracowania pokazano, że kierunek i siła zależności między niepewnością a inwestycjami mogą być różne. To, w jaki sposób niepewność oddziałuje na inwestycje, zależy od charakterystyki przedsiębiorstwa i warunków jego funkcjonowania. W celu przeanalizowania zależności będącej podstawowym problemem badawczym tej pracy skonstruowano rów-

nanie inwestycji, w którym jako główne zmienne objaśniające przyjęto oczekiwany wzrost popytu oraz miarę niepewności w obszarze popytu. Jednocześnie dołączono zmienne pozwalające na uwzględnienie uznanych w literaturze determinant siły i kierunku zależności pomiędzy niepewnością a inwestycjami. Są nimi **stopień odwracalności majątku**, w jaki inwestuje dana firma (czy jest on łatwo zbywalny czy też trudno lub niezbywalny) oraz **warunki konkurencji** na rynku produktów danego przedsiębiorstwa⁶. Do modelu dołączono również informacje na temat **ograniczenia budżetowego** badanych firm, co ma na celu sprawdzenie, jak niepewność wpływa na inwestycje poprzez ograniczone możliwości ich finansowania. Uwzględniono również zróżnicowanie wielkości badanych podmiotów, umieszczając w estymowanym równaniu inwestycji informację o wielkości zatrudnienia.

3.1. Dane oraz opis zmiennych

W prezentowanym badaniu posłużono się danymi jednostkowymi pochodzącymi z ankiety przeprowadzonej w 2006 r. przez Narodowy Bank Polski na grupie ponad siedmiuset przedsiębiorstw niefinansowych⁷. Treść pytań, które służyły do pozyskania informacji wykorzystanych w niniejszej pracy, zamieszczono w ramce 1 znajdującej się w aneksie. Ramka 3 w aneksie prezentuje strukturę badanej próby pod względem form własności i rodzaju działalności (sekcji PKD). Na podstawie uzyskanych poprzez ankietę informacji zbudowano zmienne, scharakteryzowane poniżej.

Zmienna objaśniana

Plany inwestycyjne/inwestycje – zmienna binarna określająca, czy dane przedsiębiorstwo planowało przeprowadzenie w 2006 r. poważniejszych inwestycji (istotnych z punktu widzenia dalszej działalności podmiotu i(lub) o znacznym nakładzie finansowym) inwestycji. Za przedsiębiorstwa planujące istotne inwestycje uznano podmioty, które zaznaczyły wariant A i (lub) B w pytaniu o intensywność planowanych procesów inwestycyjnych.

Zmienne objaśniające

Inwestycje realizowane – zmienna binarna określająca, czy dany podmiot przeprowadził w 2005 r. poważne inwestycje (istotne z punktu widzenia dalszej działalności podmiotu i (lub) o znacznym nakładzie finansowym).

Miara niepewności – zmienną kwantyfikującą niepewność jest miara dyspersji (odchylenie standardowe) subiektywnych (deklarowanych przez poszczególne przedsiębiorstwa) rozkładów prawdopodobieństwa oczekiwanych zmian poziomu sprzedaży produktów przedsiębiorstwa w 2006 r. Subiektywne rozkłady oczekiwanych zmian sprzedaży skonstruowano na podstawie odpowiedzi przedsiębiorstw na pytanie ankietowe.

Oczekiwany wzrost sprzedaży określono na podstawie tego samego pytania, które posłużyło do konstrukcji miary niepewności. Na podstawie subiektywnych rozkładów prawdopodobieństwa oszacowano średni oczekiwany wzrost wielkości sprzedaży (wyrażony w punktach procentowych).

Odwracalność inwestycji wyznaczono na podstawie oceny możliwości ewentualnej odsprzedaży na rynku wtórnym maszyn i urządzeń wykorzystywanych do produkcji głównego produktu. Dla każdego przedsiębiorstwa określono, czy elementy majątku, w który inwestuje dany podmiot, są łatwo zbywalne (za satysfakcjonującą cenę), czy trudno zbywalne lub niezbywalne. Inwestycje uznano za odwracalne w przypadku podmiotów, które zaznaczyły wariant A w pytaniu o możliwość odsprzedaży maszyn i urządzeń wykorzystywanych do wytwarzania głównego produktu.

Pozycja rynkowa przedsiębiorstwa została określona jako silna w przypadku, gdy dany podmiot był ważnym producentem (handlowcem) w kraju i(lub) na rynkach europejskich, czyli odpowiedział twierdząco na warianty A, B lub (oraz) C pytania zawartego w aneksie. Z kolei podmioty, które były jedną z wielu podobnych firm (zaznaczyły odpowiedź D na to samo pytanie), uznano za firmy o słabszej pozycji rynkowej działające w warunkach silnej konkurencji.

Ograniczanie finansowe oceniono na podstawie pytania o dostępność finansowania za pomocą kredytu. Za przedsiębiorstwo ograniczone finansowo uznano podmioty, które spotkały się z odmową udzielenia kredytu, czyli zaznaczyły warianty B, C i (lub) D w pytaniu zamieszczonym w aneksie, a także przedsiębiorstwa nieubiegające się o kredyt (zaznaczyły odpowiedź E), gdyż uznały, że nie mają zdolności kredytowej.

Pozostałe zmienne – do modelu dołączono również zmienną określającą wielkość przedsiębiorstwa (przybliżoną przez logarytm wielkości zatrudnienia). Zbiór zmiennych binarnych różnicujących przedsiębiorstwa ze względu na formę własności, sekcję PKD oraz lokalizację na terenie kraju (ze względu na województwo) okazał się nieistotny w objaśnianiu planów inwestycyjnych analizowanych podmiotów i nie znalazł się w estymowanym równaniu.

⁶ Siła konkurencji jest przybliżeniem elastyczności cenowej funkcji popytu, a tym samym krańcowej zyskowności inwestycji. W warunkach doskonałej konkurencji, tzn. gdy występuje pełna elastyczność cenowa popytu, zysk przedsiębiorstwa jest liniową funkcją poziomu kapitału. W przypadku gdy konkurencja jest ograniczona (brak pełnej elastyczności cenowej popytu), zysk przedsiębiorstwa przypadający na każdą kolejną jednostkę zainwestowanego kapitału jest coraz mniejszy (porównaj Caballero (1991)).

⁷ W podobnych badaniach wykorzystuje się również inne dodatkowe zbiory danych jednostkowych. Jednak ze względu na brak możliwości skojarzenia zbiorów danych zawierających wyniki ankietowe ze zbiorami zawierającymi informacje ze sprawozdań finansowych GUS (F01/F02) takie rozwiązanie nie było tutaj możliwe.

Należy podkreślić, że głównym celem modelu jest weryfikacja konkretnych hipotez badawczych. Dlatego przy doborze zmiennych objaśniających kierowano się przede wszystkim podobnymi badaniami empirycznymi i modelami teoretycznymi, na których są one oparte. Drugorzędnym celem było natomiast określenie optymalnego zbioru informacji do objaśniania prognozowanych inwestycji w badanej grupie przedsiębiorstw.

Podstawowe charakterystyki wyżej opisanych zmiennych, które zostały włączone do modelu, zaprezentowano w ramce 2 w aneksie.

3.2. Postać modelu

$$\begin{aligned} Inw^* = & \alpha_0 + \alpha_1 * INW_L + \beta_0 * SP + \beta_1 * U + \beta_2 * (U * ODW) \\ & + \beta_3 * (U * POZ) + \beta_4 * ODW + \beta_5 * POZ + \\ & \beta_6 * OGR + \beta_7 * LZATR + \varepsilon \end{aligned}$$

Ostatecznie model przyjął postać:

Ze względu na charakter zmiennej objaśnianej (dychotomicznej) do estymacji tego równania wykorzystano model probabilistyczny typu logit. Zmienną objaśnianą jest więc logarytm ilorazu szans:

$$Inw^* = \log\left(\frac{P(Inw = 1)}{P(Inw = 0)}\right),$$

gdzie $P(Inw = 1)$ to prawdopodobieństwo tego, że dana firma planowała przeprowadzenie w 2006 r. poważniejszej inwestycji, natomiast $P(Inw = 0)$ to prawdopodobieństwo zdarzenia przeciwnego.

Zmienne objaśniające:

INW_L – zmienna binarna przyjmująca wartość jeden, gdy firma realizowała (w 2005 r.) poważniejsze inwestycje, i zero w przeciwnym przypadku,

SP – oczekiwana zmiana wielkości sprzedaży, w %,

U – wskaźnik niepewności,

ODW – zmienna binarna przyjmująca wartość jeden, gdy majątek firmy jest łatwo zbywalny, i zero w przeciwnym przypadku,

POZ – zmienna binarna przyjmująca wartość jeden, gdy firma ma monopolistyczną lub silną pozycję na rynku i zero w przeciwnym przypadku,

OGR – zmienna binarna przyjmująca wartość jeden, gdy przedsiębiorstwo jest ograniczone finansowo, i zero w przeciwnym przypadku,

$LZATR$ – logarytm z wielkości zatrudnienia,

ε – reszty modelu.

Zawarte w tym równaniu interakcje pomiędzy zmienną U wyrażającą niepewność oraz odpowiednio zmiennymi POZ i ODW pozwalają na określenie kierunku i siły wpływu niepewności na inwestycje w zależności od pozycji rynkowej przedsiębiorstwa oraz odwracalności inwestycji⁸. Analiza odpowiednich sum

oszacowań parametrów β_1 oraz β_2 i β_3 informuje o zależności pomiędzy niepewnością a inwestycjami w następujących grupach przedsiębiorstw.

1. $\beta_1 + \beta_2$: pokazuje wpływ niepewności na inwestycje w przedsiębiorstwach działających w warunkach silnej konkurencji, których majątek trwały charakteryzuje się wysokim stopniem odwracalności. Taka sytuacja odpowiada założeniom modelu neoklasycznego prezentowanego w pracach Hartmana i Abła. Stwierdzili oni, że w takich warunkach wzrost niepewności przyczynia się do ożywienia inwestycji.

2. β_1 : dotyczy przedsiębiorstw działających w warunkach silnej konkurencji, których majątek trwały jest trudno zbywalny.

3. $\beta_1 + \beta_2 + \beta_3$: w przedsiębiorstwach mających silną pozycję rynkową, których inwestycje mają charakter odwracalny.

4. $\beta_1 + \beta_3$: w przedsiębiorstwach mających silną pozycję rynkową, których majątek trwały charakteryzuje się stosunkowo niskim stopniem zbywalności.

Ponieważ w niniejszym badaniu interesuje nas wyłącznie kierunek ewentualnego wpływu poszczególnych zmiennych na decyzje inwestycyjne, pominiemy interpretację wartości absolutnych oszacowań poszczególnych parametrów. Istotę analizy stanowią weryfikacja istotności poszczególnych zmiennych oraz analiza znaków przy oszacowaniach parametrów.

3.3. Wyniki estymacji

Tabela 1 prezentuje wyniki estymacji kilku wersji równania podstawowego prezentowanego w poprzednim rozdziale. Model (I) został oszacowany dla pełnej wersji równania, (II) pomija zmienne, które okazały się nieistotne w (I). W modelu (III) nie uwzględniono zmiennych wpływających na relację pomiędzy niepewnością a inwestycjami, czyli czynników różnicujących przedsiębiorstwa pod względem siły konkurencji i odwracalności inwestycji. Model (IV) prezentuje natomiast wyniki estymacji modelu szacowanego wyłącznie dla przedsiębiorstw przetwórstwa przemysłowego – najliczniejszej sekcji w próbie.

Na podstawie otrzymanych oszacowań należy stwierdzić, że inwestycje były częściej planowane przez podmioty, które realizowały istotne inwestycje rok wcześniej. Dodatkowo zaobserwowano dodatnią zależność pomiędzy wielkością oczekiwanego wzrostu popytu a prawdopodobieństwem planowania inwestycji przez przedsiębiorstwo. Wpływ ograniczenia finansowego badanych przedsiębiorstw na planowane inwestycje jest niejednoznaczny. Ograniczenie finansowe jest czynnikiem zmniejszającym planowane inwestycje w sekcji przetwórstwo przemysłowe, nie ma jednak istotnego wpływu na inwestycje na poziomie pełnej próby (parametr przy zmiennej OGR w równaniu szacowanym na

⁸ W celu minimalizacji problemu współliniowości pomiędzy zmiennymi $ODW * U$ oraz $POZ * U$ a zmienną U zmienną U zmienną U w iloczynach skorygowano o średnią.

Tabela 1. Wyniki estymacji modelu logitowego

Zmienna objaśniana Inw	Cała próba			Przetwórstwo przemysłowe
	(I)	(II)	(III)	(IV)
<i>INW_L</i>	1,427 [0,000]***	1,478 [0,000]***	1,493 [0,000]***	1,137 [0,000]***
<i>SP</i>	0,039 [0,000]***	0,039 [0,000]***	0,036 [0,000]***	0,043 [0,003]***
<i>U</i>	-0,056 [0,018]**	-0,073 [0,000]***	-0,053 [0,001]***	-0,078 [0,035]**
<i>U * ODW</i>	0,079 [0,012]**	0,068 [0,036]**		0,104 [0,020]**
<i>U * POZ</i>	-0,039 [0,194]			-0,017 [0,667]
<i>ODW</i>	-0,010 [0,972]			-0,397 [0,322]
<i>POZ</i>	0,229 [0,337]			0,294 [0,350]
<i>OGR</i>	-0,513 [0,144]			-1,294 [0,005]***
<i>LZATR</i>	0,630 [0,000]***	0,648 [0,000]***	0,657 [0,000]***	0,744 [0,000]***
Stała	-3,190 [0,000]***	-3,172 [0,000]***	-3,223 [0,000]***	-3,467 [0,000]***
Liczba obserwacji	583	583	583	335
Pseudo R ²	0,27	0,26	0,26	0,26
$\beta_1 + \beta_3 (U + U * POZ)$	-0,095 [0,001]***			-0,095 [0,006]***
$\beta_1 + \beta_2 (U + U * ODW)$	0,023 [0,393]	-0,005 [0,846]		0,026 [0,514]
$\beta_1 + \beta_2 + \beta_3$ ($U + U * POZ + U * ODW$)	-0,016 [0,565]			0,009 [0,805]

* zmienna istotna przy poziomie istotności $\alpha = 0,1$ ** zmienna istotna przy poziomie istotności $\alpha = 0,05$ *** zmienna istotna przy poziomie istotności $\alpha = 0,01$

Źródło: opracowanie własne.

całej próbie nie różni się istotnie od zera⁹). Pozycja rynkowa przedsiębiorstwa nie wpływała na intensywność planów inwestycyjnych, podobnie jak charakter majątku, w który inwestują analizowane podmioty. Jed-

⁹ Brak istotności tego parametru w dużej mierze wynika z korelacji pomiędzy inwestycjami realizowanymi w ubiegłym roku (*INW_L*) a ograniczeniem finansowym przedsiębiorstw. Gdy z modelu usunie się zmienną *INW_L*, parametr przy zmiennej *OGR* staje się istotny przy poziomie istotności $\alpha = 0,05$.

nocześnie widać, że skłonność do inwestowania zwiększa się wraz ze wzrostem wielkości przedsiębiorstwa (mierzoną przez logarytm zatrudnienia). Rezultaty te są zgodne z oczekiwaniami autorów i pozwalają na pozytywną weryfikację merytoryczną szacowanego modelu. Wyniki pokazują ponadto, że wpływ zmiennych objaśniających na plany inwestycyjne jest umiarkowany, o czym świadczy stosunkowo niska wartość współczynnika pseudo-R².

Wpływ niepewności na inwestycje oceniono na podstawie oszacowań trzech parametrów β_1 , β_2 , oraz β_3 , stojących odpowiednio przy zmiennych U , $U * ODW$ i $U * POZ$. Na podstawie modelu (III), w którym nie uwzględniono różnicowania przedsiębiorstw ze względu na odwracalność inwestycji oraz siłę konkurencji (zmiennych ODW i POZ), stwierdzono, że w analizowanej próbie wzrost niepewności wpływał negatywnie na inwestycje (na co wskazuje ujemny i istotny parametr β_1).

Jednocześnie zaobserwowano, że możliwość odsprzedaży elementów majątku (odwracalność inwestycji) zmniejsza negatywne oddziaływanie niepewności na decyzje inwestycyjne przedsiębiorstw, o czym świadczy dodatni i istotny parametr β_2 w równaniach (I), (II) oraz (IV). Wynik ten jest zgodny z zaprezentowanymi w części teoretycznej wynikami różnych modeli teoretycznych i empirycznych.

Na siłę i kierunek relacji pomiędzy niepewnością a planami inwestycyjnymi nie wpływa natomiast pozycja rynkowa badanych podmiotów (czyli zaburzenia konkurencyjnych mechanizmów rynkowych). Co prawda parametr β_3 przy interakcji zmiennych $U * POZ$ jest ujemny, co wskazywałoby, że w warunkach słabej konkurencji niepewność ma silniejszy negatywny wpływ na inwestycje. We wszystkich szacowanych równaniach owa zależność okazała się jednak statystycznie nieistotna. Tak więc ta często podkreślana w literaturze relacja nie znalazła potwierdzenia w prezentowanym powyżej modelu empirycznym.

Na dole tabeli 1 znajdują się sumy parametrów stojących w równaniu inwestycji przy zmiennej kwantyfikującej niepewność oraz interakcjach tej zmiennej ze zmiennymi ODW i POZ , a także prawdopodobieństwo, z jakim należy uznać istotność odpowiednich sum analizowanych zmiennych. Zaprezentowane wyniki wskazują jednoznacznie, że o kierunku zależności pomiędzy niepewnością a inwestycjami przesądza odwracalność majątku, w który inwestują badane podmioty. W przypadku braku odwracalności inwestycji (majątek przedsiębiorstwa jest trudno zbywalny) występuje negatywny wpływ niepewności na inwestycje. Z kolei gdy inwestycje mają charakter odwracalny, niepewność jest neutralna dla inwestycji (suma oszacowań parametrów $\beta_1 + \beta_2$ jest nieistotnie różna od 0). Powyższa teza jest prawdziwa niezależnie od tego, jaka jest pozycja konku-

rencyjna badanego przedsiębiorstwa, co sugerował już brak istotności parametru przy zmiennej $U * POZ$.

Otrzymane rezultaty pokrywają się z wynikami modelu teoretycznego Caballero w zakresie kierunku wpływu odwracalności inwestycji na relacje pomiędzy niepewnością a inwestycjami. Inaczej wygląda sytuacja z drugim czynnikiem, tzn. siłą konkurencji, który w niniejszym badaniu okazał się nieistotny dla zależności niepewność – inwestycje. Dodatkowo Caballero twierdził, że sama nieodwracalność inwestycji nie implikuje negatywnego oddziaływania niepewności na inwestycje, czego nie potwierdziły prezentowane w niniejszym badaniu wyniki analizy danych empirycznych.

4. Zakończenie

Literatura prezentująca zarówno modele teoretyczne, jak i wyniki badań empirycznych nie daje jednoznacznej odpowiedzi w kwestii kierunku wpływu niepewności na inwestycje. Zależy on bowiem od kilku czynników wymienionych we wcześniejszych rozważaniach. W niniejszym opracowaniu udowodniono, że charakter inwestycji (świadczący o symetrii lub asymetrii kosztów dostosowań) wpływa na siłę i kierunek oddziaływania niepewności na inwestycje. Nie potwierdzono natomiast wpływu pozycji rynkowej przedsiębiorstwa (będącej przybliżeniem elastyczności cenowej funkcji popytu) na siłę i kierunek relacji pomiędzy niepewnością a inwestycjami.

W badanej próbie przedsiębiorstw czynnikiem determinującym negatywny wpływ niepewności na inwestycje okazała się odwracalność inwestycji. W przypadku przedsiębiorstw inwestujących w majątek charakteryzujący się wysokim stopniem odwracalności wzrost niepewności nie miał wpływu na inwestycje, podczas gdy w podmiotach oceniających inwestycje jako nieodwracalne wpływ niepewności na inwestycje był negatywny.

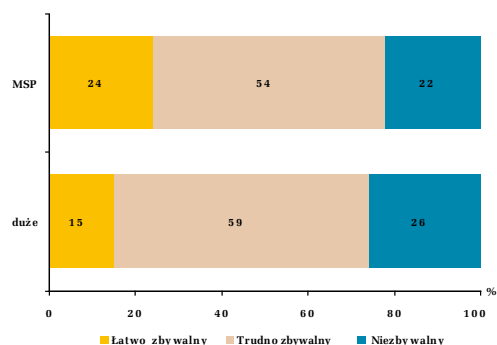
Przedstawione badanie wskazuje więc na znaczenie symetrii kosztów dostosowań poziomu kapitału dla neutralizacji negatywnego wpływu niepewności na inwestycje. Wpływ konkurencji na zmniejszenie negatywnych skutków niepewności okazał się w badanej próbie statystycznie nieistotny.

Bibliografia

- Abel A.B. (1983), *Optimal investment under uncertainty*, "American Economic Review", Vol. 73, No. 1, pp. 228–233.
- Abel A.B., Eberly J. (1994), *A Unified Model of Investment Under Uncertainty*, "American Economic Review", Vol. 84, No. 1, pp. 1369–1384.
- Bernanke B. (1983), *Irreversibility, uncertainty and cyclical investment*, "Quarterly Journal of Economics", Vol. 97, No. 1, pp. 85–106.
- Bloom N., Bond S., Reenen J. (2006), *Uncertainty and investment dynamics*, "Working Paper", No. 12383, NBER, Cambridge.
- Butzen P., Fuss C., Vermuelen P. (2002), *The impact of uncertainty on investment plans*, "Working Paper", No. 24, National Bank of Belgium, Brussels.
- Cassimon D., Engelen, P., Meersman H., Van Wouwe M. (2002), *Investment, uncertainty and irreversibility: Evidence from Belgian accounting data*, "Working Paper", No. 23, National Bank of Belgium, Brussels.
- Carruth A., Dickerson A., Henley A. (1998), *What do we know about investment under uncertainty*, "Studies in Economics", No. 9804, University of Kent, Department of Economics, Canterbury.
- Caballero R. (1991), *On the Sign of the Investment – Uncertainty Relationship*, "American Economic Review", Vol. 81, No. 1, pp. 279–288.
- Dixit A.K., Pindyck R.S. (1994), *Investment Under Uncertainty*, Princeton University Press Princeton.
- Fuss C., Vermuelen P. (2004), *Firms' investment decision in response to demand and price uncertainty*, "Working Paper", No. 347, ECB, Frankfurt.
- Ghosal V., Loughani P. (2000), *The differential impact of uncertainty on investment in small and large businesses*, "Review of Economics and Statistics", Vol. 82, No. 2, s. 338–349.
- Guiso L., Parigi G. (1999), *Investment and Demand Uncertainty*, "Quarterly Journal of Economics", Vol. 114, No. 1, pp. 185–227.
- Hartman R. (1972), *The effect of price and cost uncertainty on investment*, "Journal of Economic Theory", Vol. 5, October, pp. 258–266.
- Lerner A.P. (1934), *The Concept of Monopoly and the Measurement of Monopoly Power*, "Review of Economic Studies", Vol. 1, No. 3, pp. 157–175.
- McDonald R., Siegel D. (1986), *The value of waiting to invest*, "Quarterly Journal of Economics", Vol. 101, No. 4, pp. 707–727.
- Ninh L., Hermes N., Lanjouw G. (2003), *Irreversible Investment and Uncertainty: An Empirical Study of Rice Mills in the Mekong River Delta, Vietnam*, "Research Report", No. 03E40, SOM Research Institute, <http://irs.ub.rug.nl/ppn/260295124>
- Pattillo C. (1998), *Investment, Uncertainty, and Irreversibility in Ghana*, "IMF Staff Papers", Vol. 45, No. 3, pp. 522–553.
- Pindyck R.S. (1988), *Irreversible Investment, Capacity Choice and the Value of the Firm*, "American Economic Review", Vol. 78, No. 5, pp. 969–985.

Aneks

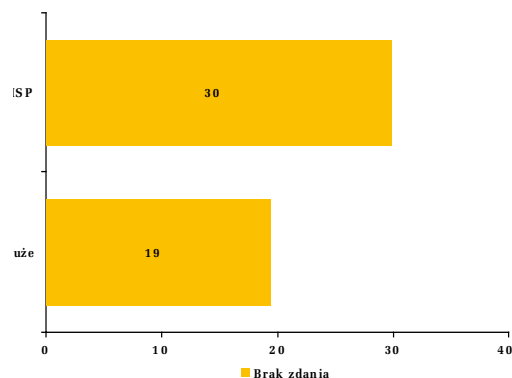
Wykres 4. Ocena odwracalności majątku wykorzystywanego do produkcji głównego produktu a wielkość przedsiębiorstwa*



* Przedsiębiorstwa przetwórstwa przemysłowego; z wyłączeniem przedsiębiorstw niepotrafiących zdecydowanie określić możliwości odsprzedaży majątku.

Źródło: badania ankietowe przedsiębiorstw, NBP.

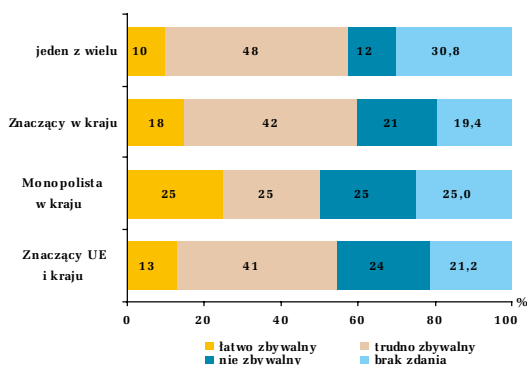
Wykres 5. Udział przedsiębiorstw niepotrafiących określić możliwości odsprzedaży majątku a wielkość zatrudnienia*



* Przedsiębiorstwa przetwórstwa przemysłowego.

Źródło: badania ankietowe przedsiębiorstw, NBP.

Wykres 6. Ocena odwracalności majątku wykorzystywanego do produkcji głównego produktu a pozycja rynkowa przedsiębiorstwa



* Przedsiębiorstwa przetwórstwa przemysłowego.

Źródło: badania ankietowe przedsiębiorstw, NBP.

Ramka 1. Pytania zadane przedsiębiorstwom w ramach rocznego badania ankietowego NBP w 2005 r., wykorzystane do konstrukcji zmiennych w modelu

Pytanie, które wykorzystano jako podstawę oceny planów inwestycyjnych

Jak przedsiębiorstwo ocenia intensywność procesów inwestycyjnych planowanych na 2006 r.?

- A. Będą to inwestycje o istotnym znaczeniu dla dalszej działalności przedsiębiorstwa
- B. Będą to inwestycje o wysokim (dużym) nakładzie finansowym z punktu widzenia przedsiębiorstwa, ale nie przełomowe dla dalszej działalności
- C. Będą to tylko nieznaczne inwestycje (małe nakłady finansowe) lub brak inwestycji
- D. Przedsiębiorstwo przewiduje wstrzymanie lub znaczne ograniczenie skali inwestycji

Analogiczne pytanie wykorzystano do oceny inwestycji realizowanych przez przedsiębiorstwo w 2005 r.

Pytanie, które wykorzystano jako podstawę oszacowania niepewności

Prosimy określić szanse realizacji każdego z poniższych scenariuszy zmian poziomu sprzedaży produktów przedsiębiorstwa w 2006 r. (XII 2006 do XII 2005, gdzie XII 2005 = 100%, przy założeniu stałych cen). W tym celu należy przyporządkować prawdopodobieństwa z przedziału pomiędzy 0–100% dla poszczególnych wariantów wydarzeń: (np.: 0% w przypadku niemożliwym do zrealizowania, 100% w sytuacji, która nastąpi na pewno, 20%, gdy dany scenariusz wystąpi z prawdopodobieństwem 20%). Suma prawdopodobieństw wystąpienia 9 scenariuszy powinna wynosić 100%.

Prognozowana zmiana sprzedaży w 2006 r.	Prawdopodobieństwo zmiany poziomu sprzedaży
wzrost powyżej 50%	
wzrost o 20–50%	
wzrost o 10–20%	
wzrost o 5–10%	
wzrost o 0–5%	
spadek o 0–5%	
spadek o 5–10%	
spadek o 10–20%	
spadek powyżej 20%	
RAZEM	100%

Pytanie, które wykorzystano jako podstawę oceny odwracalności inwestycji

Jak przedsiębiorstwo ocenia możliwości ewentualnej odsprzedaży na wtórnym rynku maszyn i urządzeń wykorzystywanych do produkcji głównego produktu? Majątek ten jest:

- A. Łatwo zbywalny (można relatywnie szybko znaleźć kupca oferującego satysfakcjonującą cenę zakupu)
- B. Trudno zbywalny (poszukiwanie kupca długotrwale i (lub) niska oferowana cena zakupu)
- C. Praktycznie niezbywalny (np. ze względu na wysoką specjalizację sprzętu)
- D. Przedsiębiorstwo nie ma orientacji na ten temat

Pytanie, które wykorzystano jako podstawę oceny pozycji rynkowej przedsiębiorstwa

Jaką pozycję zajmowało przedsiębiorstwo na rynku w 2005 r.?

- A. Było znaczącym producentem (handlowcem) na rynkach europejskich i na rynku krajowym
- B. Monopolistyczną w kraju
- C. Było znaczącym producentem (handlowcem) w kraju, ale nie odgrywało znaczącej roli na rynku europejskim
- D. Było jednym z wielu podobnych producentów (handlowców) w kraju

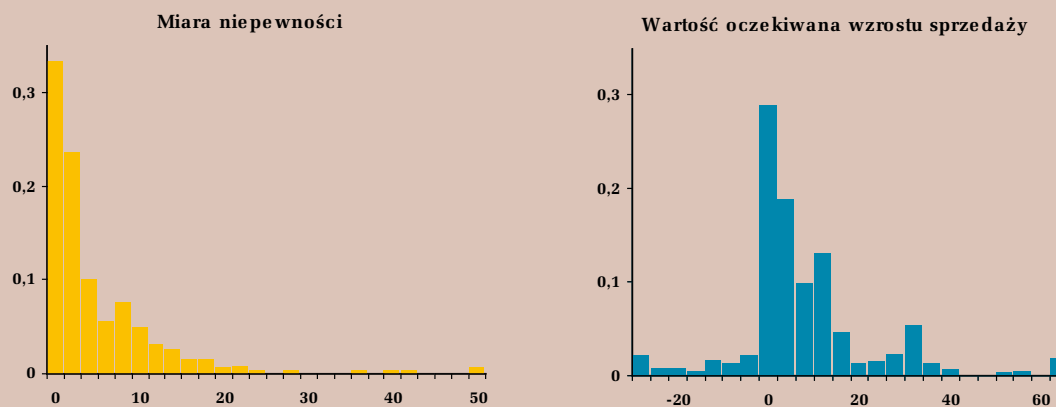
Pytanie, które wykorzystano jako podstawę oceny ograniczenia finansowego badanej firmy

Czy w 2005 r. przedsiębiorstwo spotkało się z odmową udzielenia kredytu?

- A. Nie
- B. Tak, po raz pierwszy w okresie swojego istnienia
- C. Tak, były pojedyncze przypadki odrzucenia wniosków kredytowych w okresach wcześniejszych
- D. Tak, od pewnego czasu przedsiębiorstwo często spotyka się z odmową
- E. Przedsiębiorstwo w 2005 r. nie ubiegało się o kredyt

Ramka 2 Podstawowe statystyki zmiennych włączonych do modelu

Wykres 7. Rozkład odchylenia standardowego rozkładów oczekiwanych zmian popytu (niepewności) oraz średniej oczekiwanej wielkości zmian popytu



Źródło: badanie ankietowe przedsiębiorstw, NBP

Zmienne ilościowe	Opis zmiennych	Średnia	Odch. Std.	Minimum	Maksimum	Jednostka
SP	oczekiwany wzrost sprzedaży	9,892	14,986	-30	65	pkt proc.
U	miara niepewności	5,859	9,490	0	60	pkt proc.
l_ZATR	logarytm wielkości zatrudnienia	5,356	1,410	0,69	9,95	–

Zmienne binarne	Zmienna przyjmuje wartość 1, gdy:	Częstość, z jaką zmienna przyjmuje wartość 1
INW (zmienna objaśniana)	plany inwestycyjne	0,71
INW_L	inwestycje realizowane	0,64
POZ	pozycja rynkowa	0,48
ODW	charakter majątku	0,19
OGR	ograniczenie finansowe	0,10

Ramka 3. *Struktura badanej próby (w %)*

Struktura badanej próby pod względem form własności	
Skarbu Państwa	9
Państwowych osób prawnych	5
Samorządowa	3
Prywatna krajowa	55
Zagraniczna	28
Struktura badanej próby w podziale na sekcje PKD	
Górnictwo i kopalnictwo	2
Przetwórstwo przemysłowe	58
Wytwarzanie i zaopatrywanie w energię elektryczną	6
Budownictwo	8
Handel i naprawy	13
Transport	5
Pozostałe	8

Źródło: opracowanie własne na podstawie badania ankietowego przedsiębiorstw, NBP.

Options and Market Expectations: Implied Probability Density Functions on the Polish Foreign Exchange Market*

Opcje a oczekiwania rynku: estymacja i wykorzystanie implikowanych funkcji gęstości prawdopodobieństwa na polskim rynku walutowym

Piotr Bańbula**

received: 5 June 2007, final version received: 7 April 2008, accepted: 9 April 2008

Abstract

An overview of methods used for estimation of option-implied risk-neutral probability density functions (PDFs) is presented in the study, and one of such methods, double lognormal approach, is used for the analysis of the information content of the EUR/PLN currency options on the Polish market. Estimated PDFs have proven to provide superior information concerning future volatility than historical volatility, yet their forecasting power is comparable to that of the Black-Scholes model. There are no strong grounds for using PDFs as a predictor of the future EUR/PLN exchange rate. Low informative content does not directly follow, as PDFs can be used as an indicator of markets conditions. The issues that could be addressed more thoroughly in the future studies concern the assumption of risk neutrality and the impact of the estimation method on the higher moments of the distribution.

Keywords: foreign exchange, probability density functions, option pricing, market expectations

JEL: F31, G13, D84

Streszczenie

W artykule dokonano przeglądu metod estymacji funkcji gęstości prawdopodobieństwa (PDF) instrumentu bazowego na podstawie cen opcji przy założeniu neutralności wobec ryzyka. W analizie rynku EUR/PLN zastosowano metodę dwóch rozkładów logarytmiczno-normalnych. Okazało się, że oszacowane PDF dostarczają więcej informacji o przyszłej zmienności niż zmienność historyczna, jednak ich zawartość informacyjna była bardzo zbliżona do oferowanej przez model Blacka-Scholesa. Brak jest silnych podstaw do użycia kontraktów opcyjnych jako instrumentu prognozującego przyszłe poziomy kursu EUR/PLN. Nie jest to jednak tożsame z niską zawartością informacyjną PDF, które mogą być użyte jako wskaźnik sytuacji na rynku. Elementy, które zasługują na pogłębioną analizę, to założenie o neutralności wobec ryzyka oraz wpływ metody estymacji na wyższe momenty implikowanych rozkładów prawdopodobieństwa.

Słowa kluczowe: kurs walutowy, funkcje gęstości prawdopodobieństwa, wycena opcji, oczekiwania rynku

* I am grateful to Andrzej Sławiński for drawing my attention to the topic and to Krzysztof Rybiński and Janusz Zieliński for their suggestions on the preliminary version of this article. I would also like to thank two anonymous referees for their extensive comments. Special thanks go to Łukasz Suchecki. All remaining errors are my own.

** National Bank of Poland, Domestic Operations Department; Warsaw School of Economics, e-mail: piotrbanbula@yahoo.com. The views presented in the paper are those of the author and do not necessarily reflect views of the institutions he is affiliated with.

1. Introduction

Contrary to many other financial instruments, the price of which reflects all market scenarios, options can be used to “show” probabilities attached by investors to particular events. One can obtain such information through estimation of option implied probability density functions (PDF). While this type of analysis has become increasingly popular in recent years, the number of publications in many areas of this field remains relatively limited. Most of the research has concentrated on developing, testing and comparing characteristics of new estimation techniques, whereas less attention has been paid to the analysis of their information content and forecasting power. This paper seeks to go in the latter direction and investigate the forecasting power of 1-month option contracts on the Polish foreign exchange market – namely for the EUR/PLN currency pair.

In the standard Black-Scholes option pricing model, it is assumed that the distribution of the underlying instrument is of a lognormal type.¹ Yet market prices of option contracts indicate that investors make different assumptions. These discrepancies give grounds to the analysis of options’ market quotes in order to estimate market-expected distributions of the underlying instrument, thus providing information on the expected rates of return or probability attached to particular events (realisation of a given currency level, equity price).

It should be stressed that the analysis is conducted under the assumption of risk neutrality. In some situations, such an assumption may prove to be inappropriate, resulting in bias and significant discrepancies between estimated option implied probability and subjective probability as seen by investors. Taking into consideration the unresolved difficulties in capturing agents’ preferences under different states of nature, the risk neutrality assumption dominates in works on option implied PDFs.

The paper is structured as follows. In the first part, a theoretical basis for option pricing is presented, together with a list of major anomalies between market practice and the Black-Scholes model. In particular, option prices on the market do not seem to coincide with the assumption of lognormal distribution of the underlying asset. Such observation provides basis for analysis that aims at estimating how this PDF really looks like. The second section explains how option prices can be translated into probabilities attached by market participants to certain events. The third section describes three major groups of methods that are used to estimate implied PDFs. The fourth section is dedicated to a more detailed presentation of the double lognormal approach, applied for the EUR/PLN currency pair. In the last section, the information content of the EUR/PLN option implied PDF is investigated.

¹ A lognormal distribution is a distribution of a variable whose logarithm has a normal distribution.

2. Black-Scholes model

An option² gives its buyer the right to buy or sell a given underlying instrument at the expiry date at the previously set price (strike). An option is an asymmetrical instrument as its seller has the obligation to execute a transaction on buyer’s demand at the expiry.

Valuation of each derivative requires identification of the stochastic process that governs the evolution of the underlying instrument’s price. Initially, Bachelier (1900) in his work on pricing bond option assumed that bond prices evolved according to the arithmetic Brownian motion. Still long horizon of expiry of some bonds allowed for option prices to reach negative values, which distorted valuation. An assumption of the geometric Brownian motion introduced by Samuelson (1965) for equity valuation eliminated this inconveniency. Yet his model required estimation of two important factors: the expected rate of return on equities and the discount rate. As both factors depended on investors’ utility functions, such a method of valuation suffered from the necessity of making a strong assumption on the above parameters.

Black and Scholes (1973) developed a completely novel approach to this subject. Their starting point was to analyse the transaction from the point of view of the option seller who wishes to hedge his position. As they have shown, along with an increase in hedging frequency the cost of hedging itself becomes increasingly easier to anticipate. In limiting case, where hedging activity is continuous, its cost is independent from the price of the underlying asset. The single factor that influences this cost is the variance (volatility) of the asset price. Should this variable be known in advance, it would allow for fair option valuation.

In the Black-Scholes (1973) model, it is assumed that the price of the underlying instrument evolves according to the geometric Brownian motion, which means that asset’s price can be characterised by a lognormal distribution with a constant variance. It is further assumed that the risk free rate is constant till the maturity of the contract, investors can lend and borrow at this rate and there are no transaction costs. We shall try to touch upon these assumptions later in the text.

Garman and Kohlhagen (1983) have adopted the Black-Scholes (B-S) model to currency options. Taking a similar assumption, they have shown that the prices of call (c) and put (p) options can be given by the following formulas:³

² We constrain our analysis to the European option which can be executed only at the expiry date, contrary to the American option.

³ The Garman-Kohlhagen (G-K) model does not account for discrepancies between the stochastic time (between transaction date and expiry date) and swap time (between premium being paid and final settlement of the contract). Stochastic time is linked to implied volatility of the underlying asset and swap time is linked to interest rates. The failure to account for the difference between the two measures results in incorrect option valuation. For further reading, see Stopczyński, Węgleńska (1999). Taking into consideration that possible bias due to this phenomenon is most probably significantly lower than the bias due to quality of the data, we will proceed with the analysis without correcting the equations of the G-K model.

$$c = e^{-r_f \tau} SN(d_1) - e^{-r_f \tau} XN(d_2) \quad (1)$$

$$p = e^{-r_f \tau} XN(-d_2) - e^{-r_f \tau} SN(-d_1) \quad (2)$$

where:

$$d_1 = \frac{\ln(S/X) + (r - r_f + \sigma^2/2)\tau}{\sigma\sqrt{\tau}} \quad (3)$$

$$d_2 = d_1 - \sigma\sqrt{\tau} \quad (4)$$

r and r_f – domestic and foreign risk-free rates,

τ – time to expiry (in years),

σ – standard deviation (volatility) of the underlying instrument,

S – spot foreign exchange,

X – strike price,

$N(d)$ – standard normal distribution $N(0, 1)$.

From the above equations one may see that the only unknown variable at the time when option is being priced is the volatility of the underlying instrument.⁴ That is why market makers, especially on the foreign exchange market, do not quote price at which they are ready to buy or sell an option, but the (implied) volatility. This volatility can be used to calculate option's price and thus premium to be paid via B-S model.⁵ What is important to note is that market participants do not have to "believe" in the B-S assumption to use it, as the model serves as a clear-cut transformation from volatility to prices. In other words, market makers quoting volatility obtain unequivocal information on prices they are to be paid for selling or buying a given option. Moreover, it should be noted that option-implied volatility reflects price offsetting demand with supply for options; it does not necessarily have to be equal to the expected volatility.

2.1. Anomalies in option prices – smiles and smirks

The Black-Scholes model (and its extensions) is the fundamental method for option pricing on the market. Still, market participants do alter some of the assumptions of the model that result as "anomalies" – discrepancies between B-S model implications and market quotes.

One trait of many assets' prices, including foreign exchange, is that the process governing their evolution is not of continuous nature⁶ (Micu 2005). Foreign exchange dynamics following the publication of important data may provide an example of such discontinuous adjustments. The speed and scale of price changes can

be so big that they make it virtually impossible for the option seller to hedge his exposure on continuous terms. This factor, together with transaction costs, does not allow market makers to hedge their position as effectively as it is assumed in the B-S model. To mitigate these effects, an additional premium must be included in option prices, which translates into higher quotes of implied volatility. Still, alternative valuation models that extend the B-S model to account for the above factors, such as the stochastic volatility model (Hull, White 1987; Heston 1993) or models based on jump-diffusion process (Merton 1976; Bates 1991; 1996a; 1996b), have also failed to mimic option prices on the market (Bates 2000).

The main difference between the B-S model and market practice concerns the shape of the volatility surface. Volatility surface can be generated through combining implied volatilities for different option maturities (term structure) with implied volatility for different strike prices. Information about volatility surface allows for direct valuation of any European-style option and through PDF analysis to estimate probabilities of various market scenarios in many time horizons. The B-S model implies that the volatility is equal across all strike prices and across all time horizons – thus, the volatility surface is completely flat. Even without the assumption of continuous price generating process and non-existing transaction costs, the volatility surface would tend to be higher for all maturities and all strikes in a similar scale. In reality, this does not happen and there are two main groups of anomalies between market-implied volatility surface and that implied by the B-S model. These are:

- volatility smile
- volatility term structure

Volatility term structure takes its name after the fact that implied volatility (market quotes) differs across time horizon, contrary to what the geometric Brownian motion assumption of the B-S model implies. The stylised fact on developed markets is that implied volatility rises with options' maturity, which may stem from its expected increase or from the willingness to pay additional premium for being able to hedge against detrimental price changes at longer horizons (Campa, Chang 1995).

A volatility smile refers to the situation where out-of-the-money (OTM) options exhibit higher implied volatility than at-the-money options.⁷ It means that implied volatility increases along with the distance between option's strike price and forward price. A volatility smile implies that investors do not value options assuming that the stochastic process governing the price evolution of the underlying instrument is a geometric Brownian motion. Such a phenomenon

⁴ The second variable is the interest rate which does not necessarily have to stay constant till the expiry. Still, market practice is that it is assumed to be equal to the interest rate with the same maturity as the option.

⁵ We use B-S abbreviation for currency options valuation model, though it is the G-K model that is applied. This is due to the fact that the G-K model can be treated as an extension of the B-S model to the currency market.

⁶ B-S model assumes that this process is continuous; there are no big, sudden price changes (jumps).

⁷ Option is called at-the-money (ATM) when the strike price is equal to the current price of the underlying instrument. Call (put) options are called out-of-the-money when the strike price is higher (lower) than the current price of the underlying instrument. For in-the-money options, the situation is symmetrical to the out-of-the-money case.

translates into implied leptokurtic distribution of the underlying asset (under risk-neutrality), so that the probabilities of extreme events are higher than in the lognormal distribution (fat tails effect).

The assumption of the lognormal distribution (implied by the geometric Brownian motion) was rejected in several studies on financial instruments' price behaviour – Hawkins et al. (1996); Sherrick et al. (1996); Jondeau, Rockinger (2000); Navatte, Villa (2000) and Bahra (2001).

The second phenomenon associated with volatility surface is a volatility smirk. A volatility smirk is characteristic for the equity market and for some emerging currency markets. It refers to the situation when OTM put option volatility differs from OTM call option volatility (both options having the same absolute delta values). In the case of the volatility smile, the implied volatility increased in symmetrical way along with the distance between the strike and forward price regardless of the direction (below or above forward). 'Smirk' means that this augment is not symmetrical and happens to be bigger in the space above or below forward price.

With a volatility smirk, the implied probability of a significant price increase is not equal to the probability of a significant price decrease (with the assumption of risk neutrality – more below). A volatility smirk manifests in non-zero prices of risk reversal contracts⁸ and implies skew in the expected distribution of the rates of return of the underlying instrument (assuming risk neutrality). For example, the implied volatility of the options that allow to sell one of the core markets currency (EUR, USD) against the emerging market currency below the current spot price (OTM option) tends to be lower than the implied volatility of the mirror call option with the strike price above spot price (also OTM option).⁹ Persistent difference between OTM call and put options may imply the so-called *peso problem*, a low-probability sudden decrease in value of the emerging currency (Micu 2005). Such a situation highlights the importance of risk-neutrality assumption. We shall come back to this topic later.

On the stock market, equity put options with strikes below current price exhibit higher implied volatility than call options with strikes above a current price (both with the same delta values). This volatility smirk

has become pronounced after the Black Monday on 19 October 1987, when US stock markets plunged over 20% in a single day. This is why Rubinstein (1994) describes the volatility smirk on the equity market as "*crash-o-phobia*". He suggests that investors may be afraid that such an event may re-occur, which leads to an increase in costs of protection against it (higher OTM put option prices). One can think of a currency crisis as of a stock crash happening on the foreign exchange market. A relatively large number of such extreme events materialising recently on emerging markets may, to some extent, explain the presence of a volatility smirk among emerging currencies.

3. Probability density function and Arrow-Debreu security

Due to their substantial information content, option-implied PDFs have become an increasingly popular target of scientific analysis. The prices of financial instruments reflect all (probability weighted) market-possible scenarios, yet in most of the cases it is extremely difficult, if not impossible, to derive probabilities attached to realisation of a particular event. The prime advantage of the option market is that through PDF analysis one can "show" or "see" the implied probability of any event. Such a type of research, which aims at derivation of market scenarios from the option market, can be found in Rubinstein (1994) or Ait-Sahalia, Lo (1995), where the probability of a stock market crash was estimated. On the same grounds, Melick and Thomas (1997) investigate the probability of the outcome of the Gulf War (war-peace) during the conflict in 1991; Leahy and Thomas (1996) conducted research on uncertainty due to the referendum over Quebec's independence in 1995; whereas Campa, Chang and Reider (1997; 1998) investigate consistency between official fluctuation bands in various currency regimes (like ERM) and option-implied market expectations.

The idea behind option implied PDF analysis is as follows. Assuming risk neutrality among investors, the difference between prices of two options having the same maturity but different strikes (agreed price of the underlying instrument in respect to which the transaction is to be settled at maturity) reflects probabilities that market participants attach to particular events materialising in the future.

A special case of an instrument, which has its price dependent on the future state, is an Arrow-Debreu security (Arrow 1964; Debreu 1959). This is a derivative where the buyer receives payoff of "1" if an underlying instrument takes a certain state at maturity and "0" otherwise. Assuming risk neutrality price of an Arrow-Debreu security for a certain state is directly proportional to the probability of this state to materialise.

⁸ Risk reversal (RR) on the foreign exchange market consists of two option contracts (long call option and short put option), both having the same absolute delta value. RR offer price is given as a difference between long call option volatility and short put option; bid price is taken as a difference between short call option and long put option. In other words, RR is a difference between volatility on the right side of the volatility smirk and the volatility on the left side of the smirk. In this work, RR is taken as a mean of the two prices. The most popular RR contract is 25D which consists of two options both having delta equal to 0.25. One can think of the delta as a probability for the option to expire in-the-money.

⁹ I.e. EUR/PLN OTM put option (strike below the current spot) has lower implied volatility than EUR/PLN OTM call option (strike over the current spot), both having the same absolute delta values.

Under further assumption of complete markets, where instruments for all possible future states are traded, it would be possible to derive the whole probability density functions of any asset (underlying instrument). The information content of such a derivative is substantial.

An Arrow-Debreu security can be replicated via options by combining them into a strategy called *butterfly spread*.¹⁰ Ross (1976) showed that having option prices for all possible future states (strikes) would allow us to recreate the PDF of the underlying instrument.

In reality, there are relatively few states (strikes) for which a liquid option market exists. In most of the cases, these are six price levels of the underlying instrument – according to market convention, these levels correspond to certain levels of option's delta and for strike equal to forward price. Delta can be roughly seen as an expected probability of an option buyer executing the transaction at maturity; its absolute value falls between 0 and 1.¹¹ In other words, for a call (put) option this is approximately the expected probability that the future price of the underlying instrument at the maturity is above (below) strike. Formally, delta is the first derivative of options' price with respect to underlying instruments' price.

Thus, the number of liquid options across the strike space (underlying instruments' price) is far from dense. Estimation techniques generally tend to optimise some conditions to find the PDF that best fits the data. Final PDF is found by interpolation between these knots and extrapolation outside them. In the next section, we are going to describe methods that are currently used for PDF estimation with a limited amount of options available and one of such methods shall be applied to the EUR/PLN exchange rate.

4. Probability density function estimation methods

Three main groups for PDF estimation can be identified (BIS 1999; Syrdal 2002):

(i) explicit assumption concerning type of stochastic process governing price evolution of the underlying instrument,

(ii) founded on Cox-Ross (1976) equation,

(iii) founded on Breeden-Litzenberger (1978) equation.

The second and third group dominates in the literature, their popularity relatively equal with neither of them strictly overwhelming the other.

Ad (i)

In this method, one first makes assumptions concerning characteristics of the stochastic process governing price

evolution of the underlying instrument and then uses market option prices to estimate the parameters of the process. Probability density function is thus given as a by-product. This approach has been used by Bates (1991; 1996a; 1996b) and Malz (1996). Still, this is one of the least popular methods, partly due to its relatively small flexibility. Making the assumption concerning the stochastic process of the underlying instrument implies strong restrictions on the type of the PDF. This is due to the fact that a given stochastic process cannot have many types of distribution, it has only one. The advantage of this approach over other methods is that once a stochastic process is identified, one can use this method to replicate options and hedge his exposure. It is worth mentioning that valuation models that assumed a stochastic process different to the B-S (*stochastic volatility models* – Hull, White 1987; Heston 1993; *jump diffusion process* – Merton 1976; Bates 1991) have failed to generate option prices matching those observed on the market (Bates 2000). PDF implied from such methods may thus have serious problems with mirroring investors' expectations. If one aims at deriving market expectations, he should rather focus on methods (ii) and (iii).

Ad (ii)

The idea behind this method is to make some assumptions concerning the type of the distribution (PDF), and not the stochastic process itself. The advantage of such an approach is that while a stochastic process has only one corresponding distribution, a single distribution can be generated by different stochastic processes. This makes method (ii) less restrictive than (i).

Cox and Ross (1976) have shown that assuming risk neutrality option's price can be expressed as an expected value of its future values discounted with risk-free rate. Moreover, the option's expected value depends on the probability density function. Formally:

$$c = e^{-r\tau} \int_X^{\infty} q(S_T, \gamma) (S_T - X) dS_T \quad (5)$$

$$p = e^{-r\tau} \int_0^X q(S_T, \gamma) (X - S_T) dS_T \quad (6)$$

where $c(p)$ is the price of the call (put) option, X is the strike price, τ is time to maturity, and $q(S_T, \gamma)$ is the PDF of the underlying instrument S_T at the maturity T which is characterised by the vector of the parameters.

Should infinitely dense strike space was available, there would exist only one PDF matching the data. Given the fact that only few market prices are at hand, techniques based on C-R equation generally make some assumptions concerning the characteristics of the risk-neutral PDF $q(S_T, \gamma)$, whose parameters γ (such as mean and standard deviation) are estimated to minimise the difference between theoretical option prices implied by the equations above and market prices.

¹⁰ This combination consists of a sale of two call options with strike X and purchase of two call options with strikes $X + \Delta x$ and $X - \Delta x$. In practice, *butterfly spread* is constructed out of *straddle* and *strangle* strategies – two call options (ATM and OTM) and two put options (ATM and OTM as well).

¹¹ For some exotic options it can take higher values.

One of the most popular approaches within this method is the so-called *mixture of lognormals*, applied, among others, by Melick and Thomas (1997); Mizrach (1996); Söderlind and Svensson (1997). PDF is generated by a certain number (most commonly two, hence the name *double lognormal*) of independent lognormal distributions. The final shape of the PDF is sufficiently flexible to accommodate such effects like pronounced skewness or fat tails. We shall tackle this approach more in depth in the sections to follow.

Researchers working on the implied PDF have also turned to more general types of distribution that are even less restrictive. These include g and h distribution (Dutta, Babel 2005), second type beta distribution (McDonald, Bookstaber 1991; Bookstaber, McDonald 1987; McDonald 1996; Rebonato 1999), Burr III distribution (Sherick et al. 1996) and Weibull distribution (Savickas 2001).

Works of Madan and Milne (1994) as well as Corrado and Su (1998) provide an example of a different approach, where as a starting point for estimation normal distribution is taken and is further corrected via Hermite polynomial. Rubinstein (1994) starts with a lognormal distribution and accounts for bid-ask spread in option prices. Among many types of distribution, he seeks one that corresponds to the bid-ask restrictions and is most similar to the initial lognormal distribution. Buchen and Kelley (1996) go in a similar direction, yet for optimisation procedure they use maximum entropy instead of least squares.

Ad (iii)

A connection between option prices and PDF was formally provided by Breeden and Litzberger (1978). They have shown¹² that the second derivative of options' price with respect to strike is directly proportional to the underlying instrument's probability density function and under the assumption of risk-neutrality the following relation exists:

$$\frac{\partial^2 c(X, \tau)}{\partial X^2} = e^{-r\tau} q(S_T) \quad (7)$$

where c is options' price, X corresponds to strike, τ is time to maturity and $q(S_T)$ is the PDF of the underlying instrument S_T . The big advantage of this approach is that it does not make any assumptions concerning underlying instruments' price dynamics. Methods based on the Breeden-Litzberger (B-L) equation estimate parameters of the function that binds call option prices with the strike, so that after differentiating twice, PDF can be obtained.¹³ While methods based on the Cox-Ross (C-R) equation were directly selecting (optimising) among the possible PDFs, methods based on the B-L

formula estimate the function relating option prices with strikes as a first step and then differentiate this function to obtain PDF. Moreover, in some of the cases, PDF tails must be fitted separately.

Shimko (1993) was among the very first to apply this approach. He first uses the B-S formula to change the *strike price-option price* space into the *strike price-implied volatility* space. Based on market quotes of volatility, a quadratic function relating option prices with volatility using the B-S model is estimated. Once this is achieved, one can express implied volatility as a function of the strike price. Differentiating twice, along the B-L equation, PDF is obtained. Still, the tails of the distribution that lack data must be estimated separately. Shimko (1993) uses the lognormal distribution that is fitted into tails so that the integral on the whole PDF sums up to one. Malz (1997) follows a similar path, yet in place of strikes he uses option's deltas, which allows for a better fit of the PDF. Campa, Cheng and Reider (1998) who also base on the Shimko's approach use splines instead of a quadratic equation. This provides them with greater estimation flexibility, leading to better PDF fit. Bliss and Panigirtzoglou (2000) combine Malz's (1997) approach of taking volatility at certain delta values with spline estimation adopted by Campa, Cheng and Reider (1998).

Ait-Sahalia and Lo (1998) take quite a different path. Instead of estimating the relation between options' price and strike at each point of time separately, as was done by other researchers, they estimate the general form of the function for the whole time series available. To this end, Naradya-Watsan kernel estimator is applied, yet this approach tends to be extremely time consuming.

It must be stressed that all of the above mentioned methods and approaches of PDF estimation implicitly or explicitly assume risk-neutrality among investors. Taking option prices from the risk-averse world to estimate PDF in the risk-neutral model can result in a serious bias. In such a case, an event (state) with small occurrence probability in the world characterised by risk aversion may have substantially higher probability attached to it in the world assuming risk-neutrality. This is due to the fact that risk aversion, or even loss aversion, implies that certain states of nature will have a relatively high value attached to them, not necessarily reflecting true probability of occurrence, but simply people's willingness to pay substantially more for the insurance against them than the actuarial cost. The final price of the instrument will thus depend on two factors – subjective probability for the state to materialise and utility attached to it. While operating in risk-neutral model we shall not be able to distinguish between the two, which may lead to overstatement of certain probabilities. Only recently some researches have started to tackle risk-aversion in PDF estimation

¹² Assuming there are no transaction costs, no restrictions on short sale and investors can lend and borrow at risk-free rate.

¹³ After single differentiation one can obtain cumulated density function.

(Aït-Sahalia, Lo 2000; Jackwerth 2000; Rosenberg, Engle 2002). Still, the problem of capturing agents' utility across different states of nature remains unresolved and the risk-neutrality assumption remains the dominant approach in PDF estimation studies. Bearing this in mind, one should be careful when drawing conclusions concerning market expectations derived from option implied risk-neutral PDFs.

4.1. Methods comparison

Taking into account the number of approaches towards option implied risk-neutral PDF estimation, one would most certainly be interested in their performance, strengths and weaknesses. A comparison between various methods can be found in Campa et al. (1998); Coutant et al. (2000); Mc Manus (1999); Jodeau, Rockinger (2000); Syrdal (2002) and Dutta, Babel (2005). The most common criteria include goodness of fit and solution stability. The first approach is generally based on comparison between market prices and theoretical (model-implied) prices. To this end, various statistics are applied (MSE, MSPE, MAE, MAPE, RMSE).

These analyses indicate that while each method has its strengths and weaknesses, the differences in the first two moments of the implied PDF (mean and standard deviation) tend to be very small (Jackwerth 1999; Micu 2005). Should a researcher be interested in these parameters he can choose relatively freely among possible methods without significant losses in precision of the estimation.

However, higher moments of the implied distribution (skewness and kurtosis) do tend to be prone to estimation method (Mc Manus 1999). The differences among them stem from relatively high dependence between estimation procedure and implied probabilities outside the 10th and 90th percentile (Melick 1999). These extreme probabilities have in turn a significant influence on higher moments of the distribution. As reliable market data for extreme OTM options are very difficult to encounter, there are no strong grounds for deciding which method is most accurate in estimating distribution's higher moments (Cooper 1999). Taking into consideration that one cannot be sure about the extent to which implied skewness and kurtosis are influenced by the given procedure, the information content of option implied risk-neutral PDF is diminished.

In the analysis of the EUR/PLN exchange rate, the double lognormal (DLN) method is applied. The decision was based on the following criteria:

- DLN is one the most popular approaches
- some analyses (Mc Manus 1999, Syrdal 2002) indicate that while the differences between various methods in terms of goodness of fit tend to be small, the DLN approach has the best performance
- DLN is relatively easy to implement

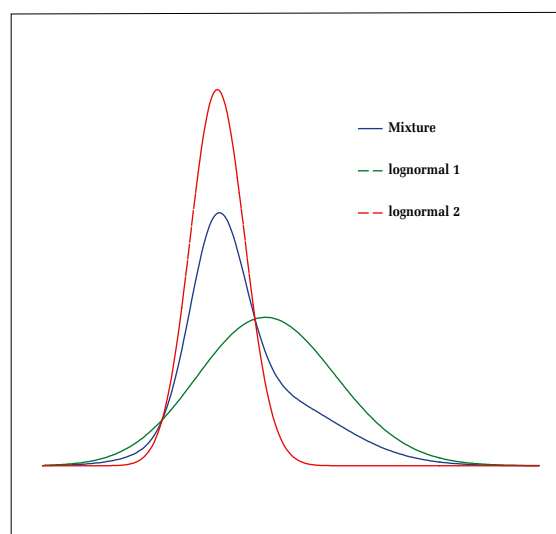
One serious drawback of the DLN method is its instability in the environment of low volatility and pronounced skewness (Cooper 1999). This problem was partly solved by imposing additional restrictions on the model (following Syrdal 2002).

5. The double lognormal approach

The method that mixes an independent lognormal distribution to obtain PDF was popularised by Melick and Thomas (1997) who applied it for PDF estimation on the oil options market. In their study, the final option implied risk-neutral PDF was a mixture of three LN distributions, still subsequent analyses used only two distributions (Bahra 1997; Söderlind, Svensson 1997; McManus 1999; Syrdal 2002; Micu 2005). The argument for such an approach is twofold. First, two LN distributions do guarantee flexibility concerning the shape of the final PDF, so that high goodness of fit is maintained. Secondly, DLN is free from problems occurring when one wants to estimate a relatively big number of parameters having limited number of option prices – as was in the original approach.

Let us recall that the DLN belongs to the second (ii) group of methods which build on the observation made by Cox and Ross (1976) that options prices under risk neutrality depend on the probability density function of the underlying asset. The idea behind these methods is to make some assumptions concerning the possible type of the final risk-neutral PDF, $q(S_T, \gamma)$, and then to estimate its parameters γ

Figure 1. Visualisation of the double lognormal approach in the rates of return domain



Source: Own calculations.

so that the difference between theoretical prices and market prices is minimised. Within the DLN approach, the final PDF is obtained by mixing two lognormal distributions with five parameters being estimated. These parameters are mean and standard deviation for each independent lognormal distribution (α_1, β_1 and α_2, β_2 accordingly) and weight θ attached to one of the distribution, with sum of weights equal to 1.¹⁴ Formally, final PDF is given as:

$$q(S_T) = \theta \cdot L(S_T | \alpha_1, \beta_1) + (1 - \theta) \cdot L(S_T | \alpha_2, \beta_2) \quad (8)$$

Bahra (1997) shows that should the PDF be a product of two lognormal distributions, theoretical call (c) and put (p) option prices can be expressed analytically¹⁵:

$$c = e^{-r\tau} \left\{ \theta \left[e^{\frac{\alpha_1 + \frac{1}{2}\beta_1^2}{2}} N(d_1) - XN(d_2) \right] + (1 - \theta) \left[e^{\frac{\alpha_2 + \frac{1}{2}\beta_2^2}{2}} N(d_3) - XN(d_4) \right] \right\} \quad (9)$$

$$p = e^{-r\tau} \left\{ \theta \left[-e^{\frac{\alpha_1 + \frac{1}{2}\beta_1^2}{2}} N(-d_1) + XN(-d_2) \right] + (1 - \theta) \left[-e^{\frac{\alpha_2 + \frac{1}{2}\beta_2^2}{2}} N(-d_3) + XN(-d_4) \right] \right\} \quad (10)$$

where:

$$d_1 = \frac{-\ln(X) + \alpha_1 + \beta_1^2}{\beta_1} \quad d_2 = d_1 - \beta_1 \quad (11)$$

$$d_3 = \frac{-\ln(X) + \alpha_2 + \beta_2^2}{\beta_2} \quad d_4 = d_3 - \beta_2 \quad (12)$$

r – domestic interest rate,

τ – time to maturity (in years),

X – strike price,

$N(d)$ – standard normal distribution $N(0, 1)$.

α, β – a mean and standard deviation of each independent L-N distribution.

Moreover, mean of the final PDF must equal to the forward exchange rate – this stems from the UIP condition and assumption of risk-neutrality. Formally:

$$\theta \cdot e^{\frac{\alpha_1 + \frac{1}{2}\beta_1^2}{2}} + (1 - \theta) \cdot e^{\frac{\alpha_2 + \frac{1}{2}\beta_2^2}{2}} = F \quad (13)$$

The DLN aims to find five unknown parameters $\alpha_1, \beta_1, \alpha_2, \beta_2$ and θ , so that the sum of squared differences between theoretical DLN prices and market-observed prices is minimised. Formally, the problem can be written as follows:

$$\min_{\alpha_1, \alpha_2, \beta_1, \beta_2, \theta} \left\{ \sum_{i=1}^k (c_i - c_i^*)^2 + \sum_{j=1}^m (p_j - p_j^*)^2 + \left[\theta \cdot e^{\frac{\alpha_1 + \frac{1}{2}\beta_1^2}{2}} + (1 - \theta) \cdot e^{\frac{\alpha_2 + \frac{1}{2}\beta_2^2}{2}} - F \right]^2 \right\} \quad (14)$$

where:

c, p – theoretical DLN call and put option prices,

c^*, p^* – market call and put option prices,

k, m – number of available market call and put option prices with different strikes.

In order to mitigate undesired effects that the DLN can produce (local, very pronounced maximums of the PDF), additional restrictions have been put on the relation between standard deviations of the two lognormal distributions (following Syrdal 2002):

$$0.25 \leq \frac{\beta_1}{\beta_2} \leq 4 \quad (15)$$

5.1. Data

Data on the implied volatility and other variables come from Bloomberg and daily observations from January 2004 to February 2007. Mid-market quotes were used. Market convention is that market makers on the foreign exchange market do not quote option prices for certain strikes but rather volatility for certain delta values. In standard situations, quotes are available for deltas corresponding to five different strike prices:

- one for zero-delta straddle strategy – a combination of call and put option with the same strike¹⁶, with the overall forward delta of zero (premium paid in base currency)

- two for OTM call option with exercise probability at approximately 10% and 25% (10-delta¹⁷ call and 25-delta call), with strike price above the forward rate,

- two for OTM put option with exercise probability at approximately 10% and 25% (10-delta put and 25-delta put), with strike below the forward rate,

One can see that we do not have information concerning volatility-strike space or option price-strike space. Still, using the B-S formula, one can relatively easily translate the volatility-delta space into volatility-strike space and then further into option price-strike space. In a first step, we apply the B-S formula to calculate strike prices corresponding to delta values for which implied volatility is quoted. Then, B-S is once again used to derive option prices from volatility quotes. At this point we do have all the necessary information to calculate theoretical DLN option prices and we solve the optimisation problem from the equation 14. All calculations were carried out in Microsoft Excel with Solver package.¹⁸ While the choice of points on the volatility smile is imposed by market convention, the positive aspect is that these points cover a relatively wide spectrum of the implied PDF. Still, the number of

¹⁶ Such a strike lies close, but it's not equal to the forward price of the underlying instrument.

¹⁷ 10-delta in market notation describes option with delta value of 0.1; call or put further informs whether one deals with buy or sell option.

¹⁸ Dutta, Babel (2005) also applied Solver package to solve their optimisation problems, with deliberate resignation from MATLAB.

¹⁴ So that the integral of the final PDF is also equal to 1.

¹⁵ One could see that the G-K model can be viewed as a special case of DLN, where, among others, $\theta = 1$.

points used in the optimisation procedure is limited to five (corresponding to approximately 10%, 25%, 50%, 75% and 90% exercise probability, with 50% being the forward rate).

As always, the quality of the data can have a significant effect on the PDF estimation. Possible problems stem from (Melick, Thomas 1997; Bliss, Panigirtzoglou 1999; Syrdal 2002):

- liquidity (probably lower for deep OTM options)
- bid-ask spread in prices of both options and underlying instrument
- not necessarily simultaneous quotes for option and underlying instrument
- errors in data of information systems
- inclusion of option prices that were not traded but only quoted.

Checking the no-arbitrage conditions eliminated some of the most obvious errors. Still, it is probable that the remaining noise in the data on the Polish market significantly exceeds that encountered on the most developed markets.

5.2. Goodness of fit

As Mean Average Percentage Error (MAPE) calculated on the whole sample does not exceed 0.75%¹⁹, model's goodness of fit can be thought as quite high. Thus, the main problems one can encounter while estimating the implied PDF stem rather from the quality of the data and risk-neutrality assumption. It seems that the assumption of risk neutrality can be of considerable importance especially on the emerging markets, i.e. due to *peso problem*, which makes them more similar to stock markets in respect to persistent skew in the implied PDF.

6. Information content of the option implied risk-neutral probability density function

There are basically two types of studies on the information content of option-implied risk-neutral PDF.

¹⁹ Calculated with the exclusion of the forward rate, which is a more restrictive approach (forward rate tends to be many times higher than the remaining four OTM option prices)

The first group uses PDF to analyse the probabilities of abrupt changes in asset prices or, more generally, market expectations, not necessarily checking for their forecasting power. Examples of such studies have been presented earlier in this paper (Rubinstein 1994; Aït-Sahalia, Lo 1995; Melick, Thomas 1997; Leahy, Thomas 1996; Campa et al. 1997; 1998).

The second group consists of studies that verify the forecasting power of the option implied risk-neutral PDF. The main area of research within this approach concerns PDF's predictive abilities over future volatility of the underlying asset. To this end, linear regression is commonly applied (Canina, Figlewski 1993):

$$\sigma_{real} = \alpha + \beta\sigma_k + \varepsilon_t \quad (16)$$

where σ_{real} describes realised (future) volatility, and σ_k corresponds to different measures of volatility used for forecasting. These include ATMF implied volatility (basically the B-S formula), PDF implied volatility or historical (rolling) volatility. In accordance with market convention, all volatility measures are annualised. ATMF implied volatility is directly taken from Bloomberg, whereas PDF implied volatility is calculated with a formula advocated by Jarrow and Rudd (1982):

$$\sigma_{PDF} = \sqrt{\ln\left(\frac{\sigma^2}{\mu^2} + 1\right) / \tau} \quad (17)$$

where σ and μ correspond to PDF's standard deviation and mean, τ is time to maturity in years. Historical (past, rolling) 1 month volatility observed a day before the contract was sold is calculated in a standard way as a standard deviation (s) of daily foreign exchange changes divided by the square root over time to maturity:

$$\sigma_{HIST} = \frac{s}{\sqrt{\tau}} = \sqrt{\frac{1}{\tau} \frac{1}{n-1} \sum_{i=1}^n (R_i - \bar{R})^2} \quad (18)$$

τ^* is time to maturity in years calculated with respect to working days (it is assumed that there are 252 working days within a year); by annualisation we obtain $\tau^* = 1/252$. R corresponds to a daily rate of return assuming continuous capitalisation:

$$R = \ln\left(\frac{EURPLN_{t+1}}{EURPLN_t}\right) \quad (19)$$

Table 1. Volatility forecasting

$\sigma_{real} = \alpha + \beta\sigma_k + \varepsilon_t$				
σ_k	α	β	R^2	(p-value) for β
PDF	0.0163 (0.0093)	0.6875 (0.0972)	0.207	(0.000)
ATMF	0.0206 (0.009)	0.6856 (0.0970)	0.200	(0.000)
Historical 1M	0.0627 (0.0061)	0.237 (0.0635)	0.059	(0.000)

Standard errors corrected for overlapping sample using Newey-West (1987) procedure, lag truncation = 6. Asymptotic Newey-West standard errors in parentheses.

Source: Own calculations.

To apply the above regression model to the available data set, one has to take into account the overlapping sample problem. As the frequency of the observation (daily) is higher than the frequency of the data (contracts with monthly maturity), error term in the equation is subject to autocorrelation. Still, one can suspect that after some lag autocorrelation between error terms should be close to zero. Therefore, standard errors are corrected for serial correlation using the Newey-West (1987) procedure. The computations were carried out in EViews.

The results (see Table 1) indicate that the B-S implied volatility (ATMF) and PDF implied volatility provide statistically significant information concerning future volatility. Both measures are highly correlated (see Appendix) and are characterised by very similar forecasting power. While the information content of the historical (past) 1M volatility is also significant, its predictive power is considerably lower. These results do coincide with the results for more developed markets in virtually all aspects mentioned, even with regard to coefficient values (Jorion 1995; Christensen, Prabhala 1998; Weinberg 2001).

While the number of studies on the forecasting power of option implied risk-neutral PDF with respect to future volatility is relatively small, the studies on higher moments of the distribution (skewness or kurtosis) are much more scarce. The majority of works on PDFs concentrates on estimation methods and comparisons between them. Weinberg (2001) indicates that the relatively small number of studies on the PDF information content may stem from the fact that the risk-neutrality assumption does not need to correspond to real world situation. Still, should this assumption generate such a serious bias, what makes researches work on the seemingly pointless task? While this question remains without satisfactory answer, one would expect that risk-neutrality assumption could result in an estimation bias. Even though the risk-neutrality assumption accompanied many studies on the bond market, the impact of this factor seems to differ across asset classes, and could be pronounced when tail, low-probability events on the emerging or equity markets are concerned.

The BIS (1999) review indicates that one of the problems with the assessment of the PDF's information content is that PDF represents a whole range of probabilities, whereas in the end only one of these events materialises. So, as long as the realised event has non-zero occurrence probability in the previously estimated PDF, one cannot reject the viability of the PDF. This argument seems to be flawed. One could still study how the *series* of the realised asset prices corresponds to the *series* of the implied PDF and on this ground analyse whether PDF is a good gauge of real distribution or it is somehow biased. To this end, the statistics based on the empirical distribution function (Stephens 1974) can be

used. Due to the fact that such a type of analysis requires relatively long series without overlapping observations, it cannot be implemented in this study.

Another obstacle for the studies on distributions' higher moments is due to the differences between methods and approaches in tail probability estimations. Tail of the distribution, where usually no data is available, may have strong impact on the skewness and kurtosis. As for now, it is hard to say which method is advisable to capture these extreme probabilities (Cooper 1999).

Bearing in mind the above mentioned problems, we will still try to assess the information content of the implied skew, foremost in terms of forecasting future exchange rate in a 1-month perspective. Let us recall that the skew in the implied distribution stems from non-zero prices of the risk-reversal (RR) contracts. RR for the EUR/PLN is persistently positive, which implies positive skew in the implied distribution. Assuming risk-neutrality, this means that the implied probability of a significant weakening of the zloty is greater than the probability of its significant strengthening. At the same time, positive skew implies that the probability mass is concentrated below the forward rate, corresponding to the higher probability of appreciation with respect to distribution's expected value. Thus, the median of the distribution lies below the forward rate. One can say that positive RR implies higher probabilities of zloty's appreciation with respect to the forward rate and, at the same time, higher probability attached to zloty's significant depreciation than to its significant appreciation. In the period under analysis, the zloty generally remained in the appreciation trend. This combination of appreciation and positive implied skew is consistent with the results obtained for other currency pairs (Malz 1997; Weinberg 2001; Campa, Chang and Reider 1998). Considering permanently positive skew in the implied PDF, one would perhaps think about the analysis that focuses on the changes in the skew.

Available studies on the predictive power of the implied skew (Weinberg 2001; Syrdal 2002) conclude that it is very small. While analysing the problem, Weinberg (2001) hints that: (i) it is difficult to say whether lack of predictive power implies low information content, (ii) risk-neutrality assumption may bias the estimation, (iii) forecasting future skew may simply be an uneasy task.

PDF estimation procedure may have a significant, possibly undesirable, impact on the distribution's tails, and consequently on the implied skew. To somehow mitigate these effects, the Pearson median skewness is used. Its advantage is that it is less sensitive to values at the ends of the distribution. On the other hand, it is not directly comparable with other skewness measures (see below). The Pearson median skewness is given as:

$$Q = \frac{3(\text{mean} - \text{median})}{\text{standard deviation}} \quad (20)$$

Table 2. *Forecasting power of the implied skewness – sign test*

Test type	No. of $z = 1$	Accuracy	(p-value)
Standard	24	63%	(0.14)
Modified ($Q_i - Q_{\text{mean}}$)	22	57%	(0.42)

Source: Own calculations.

Two types of tests are used for the analysis of the information content of the implied skew. The first one is a simple sign test of the following nature (Syrdal 2002):

$$z = 1 \begin{cases} \text{if } Q > 0 \text{ and } F - S > 0 \\ \text{or if } Q < 0 \text{ and } F - S < 0 \end{cases}$$

$$z = 0 \text{ in other cases}$$

where Q is the Pearson median skewness, F corresponds to the forward rate with maturity equal to that of the option contract and S is the future (realised) spot rate at maturity. The above test can be understood as follows. If the implied skew is positive and the spot rate lies below the forward rate, then, considering the location of the probability mass, the prediction has proven to be correct and to such event we attach value of 1. The statistical significance of the skew indications was verified with the EViews package, assuming that lack of forecasting power implies the median result of the test equal to 0.5. Given that the Newey-West (1987) procedure cannot accommodate such types of tests, the analysis was carried on the non-overlapping sample of 38 monthly observations. The persistently positive implied skew in the EUR/PLN may raise concerns over its unbiasedness as a predictor of market expectations. To partly accommodate the phenomena, a modified test is also carried out, where from each Q the mean in-the-sample skew is subtracted. In such an approach, a positive skew is deemed to be a normal situation and deviations from this 'equilibrium' are thought to provide some insight concerning market expectations.

Given the strong appreciation of the zloty in 2004, with no pronounced and persistent depreciations afterwards, the test is biased towards the rejection of the hypothesis of no forecasting power. Yet, as it was in the case of more developed markets, the results for the EUR/PLN are not very optimistic and the indications of the implied skew are not statistically significant.

Let us note that the sign test compared only the consistency between realised exchange rate changes and the indications stemming from the skew of the option-implied risk-neutral PDF. Therefore, there was no differentiation between small and big exchange rate changes. In other words, while the predictions with respect to the direction have proven to be only moderately accurate, perhaps the profit from investment strategy based on the implied skew might be still significant, yet not accommodated within the sign test.

Such a type of analysis – aiming at welfare perspective – is certainly more desirable. In our case, its application is not recommended. This is due to the fact that zloty's in-sample significant appreciation combined with the persistently positive implied skew would result in an even greater bias than observed in the sign test – the mean-skew adjusted test would still be flawed.

In order to make use of the whole data set available, another test on the forecasting power of the implied skew was carried out. It is based on the following regression:

$$\ln\left(\frac{\text{spot}_{t+k}}{\text{forward}_{t+k}^f}\right) = \alpha + \beta\phi_t + \varepsilon_t, \quad (21)$$

where forward_{t+k}^f is the forward rate with maturity equal to that of the option contract, spot_{t+k} is the future (realised) spot rate at maturity and ϕ_t corresponds to different measures of the skew of the option implied risk-neutral PDF. The test logic is thus somewhat similar to that of the sign test; still the magnitude of the exchange rate changes is allowed to vary along with skewness measures. The overlapping sample problem was solved by applying the Newey-West (1987) procedure. The calculations were carried out in EViews. The Pearson median skewness was kept as a basic measure of the skew, with some adjustment. The adjustment consisted of subtracting from each Q the skew of the lognormal distribution with standard deviation equal to the volatility in the B-S formula (ATMF). The idea behind is that the B-S model assumes positive skew in the distribution of the underlying instruments' price (and zero skew in rates of return). One would expect that the B-S implied skew should not have any forecasting power, so it is used to adjust the PDF implied skew. The skew of the B-S distribution is given by the standard formula for lognormal distribution:

$$\text{skew} = (e^{\sigma^2} + 2) \times \sqrt{e^{\sigma^2} - 1} \quad (22)$$

where ATMF volatility is used as a measure of standard deviation σ . Still, as it was mentioned earlier in the text, the above and Pearson's measures of skew are not directly comparable, so the test possibly includes some unknown bias. Thus, a simpler form of the test is also carried out, with the Pearson median skewness adjusted by the sample mean Q .²⁰

²⁰ Within this framework, there is virtually no difference between such an approach and one based on the unadjusted Pearson median skewness – the only thing that changes is the intercept (by the value of the mean skew).

Table 3. *Forecasting power of the implied skewness - regression*

$\ln\left(\frac{spot_{t+k}}{forward_{t+k}}\right) = \alpha + \beta\phi_t + \varepsilon_t$				
ϕ_t	α	β	R^2	(p-value) for β
SkewPearson PDF – SkewATMF	-0.0075 (0.0021)	-0.0095 (0.0151)	0.003	(0.53)
Pearson PDF	-0.0077 (0.0020)	-0.0263 (0.0199)	0.012	(0.19)

Standard errors corrected for overlapping sample using Newey-West (1987) procedure, lag truncation = 6. Asymptotic Newey-West standard errors in parentheses.

Source: Own calculation⁸.

While the β coefficients have the expected sign (based on the distribution of the probability mass), they are not statistically different from zero. Intercepts are found to be statistically significant most probably due to the pronounced and consistent appreciation of the zloty within the sample. The results coincide with earlier observations from the sign test.

Concluding, it seems that there are no grounds for using the skew of the option implied risk-neutral PDF for forecasting the EUR/PLN. It seems that the implied PDF could rather be used as an approximation of market situation, not necessarily informing about market expectations.

7. Conclusions

Due to their substantial information content, option-implied risk-neutral PDFs have become an increasingly popular target of scientific analysis. The prices of financial instruments reflect all (probability weighted) market-possible scenarios, yet in most of the cases it is extremely difficult, if not impossible, to derive probabilities attached to realisation of particular events. The prime advantage of the option market is that through PDF analysis one can approximate the implied probability of any event. Since option prices on the market do not seem to coincide with the Black-Scholes assumption of lognormal distribution of the underlying asset, an analysis aiming at estimating how the market-implied risk-neutral PDF really looks like is even more

interesting. Still, the number of liquid options across the strike space (underlying instruments' price) is far from dense. Thus, estimation techniques generally tend to optimise some conditions to find the PDF that best fits the data.

Most of the research has concentrated on developing, testing and comparing characteristics of various estimation techniques, whereas less attention has been paid to the analysis of their information content and forecasting power. This paper went in the latter direction and investigated the forecasting power of the EUR/PLN 1 month options. To this end, double lognormal approach was applied.

Volatility forecasted based on the PDF analysis has proven to provide better information concerning future volatility than historical volatility. Yet PDF implied volatility is highly correlated with the volatility in the Black-Scholes model (ATMF), with very similar forecasting power. As it is the case for more developed markets, there are no grounds for using the skew of the option-implied risk-neutral PDF for forecasting the zloty exchange rate. Low information content does not directly follow as estimated PDF could be used as an approximation of market situation, not necessarily informing about market expectations. Moreover, the quality of the data, an unknown relation between applied estimation technique and tail probabilities together with the accompanying assumption of risk neutrality suggest caution in drawing conclusions from available estimations, especially concerning higher moments of the distribution.

References

- Ait-Sahalia Y., Lo A.W. (1998), *Nonparametric estimation of state-price densities implied in financial asset prices*, "Journal of Finance", Vol. 53, April, pp. 499–547.
- Ait-Sahalia Y., Lo A.W. (2000), *Nonparametric Risk Management and Implied Risk Aversion*, "Journal of Econometrics", Vol. 94, No. 1–2, pp. 9–51.
- Arrow K.J. (1964), *The Role of Securities in the Optimal Allocation of Risk-bearing*, "Review of Economic Studies", Vol. 31, No. 2, pp. 91–96.
- Bachelier L. (1900), *Théorie de la spéculation*, "Annales scientifiques de l'Ecole Normale Supérieure", 3e série, tome 17, pp. 21–86.
- Bahra B. (1997), *Implied Risk-Neutral Probability Density Functions from Options Prices: Theory and Application*, "Working Paper, No. 66, Bank of England, London.

- Bahra B. (2001), *Implied risk-neutral probability density functions from option prices: A central bank perspective*, in: J. Knight, S. Satchell (eds.), *Forecasting volatility in the financial markets*, Butterworth-Heinemann, Oxford.
- Bates D.S. (1991), *The Crash of 87: Was It Expected? The Evidence from Options Markets*, "Journal of Finance", Vol. 46, July, pp. 1009–1044.
- Bates D.S. (1996a), *Jumps and Stochastic Volatility: Exchange Rate Processes Implicit in Deutsche Mark Options*, "Review of Financial Studies", Vol. 9, No. 1, pp. 69–107.
- Bates D.S. (1996b), *Dollar jump fears, 1984-1992: Distributional abnormalities implicit in currency futures options*, "Journal of International Money and Finance", Vol. 15, No. 1, pp. 65–93.
- Bates D.S. (2000), *Post-'87 crash fears in the S&P 500 futures option market*, "Journal of Econometrics", No. 94, pp. 181–238.
- BIS (1999), *Estimating and Interpreting Probability Density Functions*, Proceedings of the workshop held at BIS on 14 June 1999, Basle.
- Black F. Scholes M. (1973), *The Pricing of Options and Corporate Liabilities*, "Journal of Political Economy", Vol. 81, May – June, pp. 637–659.
- Bliss R.R., Panigirtzoglou N. (2002), *Testing the Stability of Implied Probability Density Functions*, "Journal of Banking and Finance", Vol. 26, pp. 381–422.
- Bookstaber R.M., McDonald J.B. (1987), *A general distribution for describing security price returns*, "Journal of Business", Vol. 60, pp. 401–424.
- Breeden D.T., Litzenberger R.H. (1978), *Prices of state-contingent claims implicit in option prices*, "Journal of Business", Vol. 51, pp. 621–651.
- Buchen P.W., Kelley M. (1996), *The Maximum Entropy Distribution of an Asset Inferred from Option Prices*, "Journal of Financial and Quantitative Analysis", Vol. 31, No. 1, pp. 143–159.
- Campa J.M., Chang K.P.P.H. (1995), *Testing the expectations hypothesis of the term structure of volatilities in foreign exchange options*, "Journal of Finance", Vol. 50, pp. 529–547.
- Campa J.M., Chang K.P.P.H., Reider R. (1997), *ERM Bandwidths for EMU and After: Evidence from Foreign Exchange Options*, "Economic Policy", Vol. 24, pp. 55–89.
- Campa J.M., Chang K.P.P.H., Reider R. (1998), *Implied Exchange Rate Distributions: Evidence from OTC Option Markets*, "Journal of International Money and Finance", Vol. 17, No. 1, pp. 117–160.
- Canina L., Figlewski S. (1993), *The Informational Content of Implied Volatility*, "Review of Financial Studies", Vol. 6, No. 3, pp. 659–681.
- Christensen B.J., Prabhala N.R. (1998), *The relation between implied and realized volatility*, "Journal of Financial Economics", Vol. 50, pp. 125–150.
- Cooper N. (1999), *Testing techniques for estimating implied RNDs from the prices of European-style options*, in: BIS, *Estimating and Interpreting Probability Density Functions*, Proceedings of the workshop held at BIS on 14 June 1999, Basle.
- Corrado C.J., Su T. (1996), *Skewness and Kurtosis in S&P 500 Index Returns Implied by Option Prices*, "Journal of Financial Research" Vol. 19, No. 2, pp. 175–192.
- Coutant S., Rockinger M., Jondeau E. (2001), *Reading PIBOR Futures Options Smile: The 1997 Snap Election*, "Journal of Banking and Finance", Vol. 25, No. 11, pp. 1957–1987.
- Cox J.C., Ross S.A. (1976), *The Valuation of Options for Alternative Stochastic Processes*, "Journal of Financial Economics" No. 3, pp. 145–166.
- Debreu G. (1959), *Theory of Value*, Yale University Press, New Haven and London.
- Dutta K., Babel D. (2005), *Extracting Probabilistic Information from the Prices of Interest Rate Options: Tests of Distributional Assumptions*, "Journal of Business", Vol. 78, No. 3, pp. 841–870.
- Garman M.B., Kohlhagen S.W. (1983), *Foreign currency option values*, "Journal of International Money and Finance", Vol. 2, May, pp. 231–237.
- Hawkins R. J., Rubinstein M., Daniels G.J. (1996), *Reconstruction of the probability density function implicit in option prices from incomplete and noisy data*, in: K.M. Hanson, R.N. Silver (ed.), *Maximum entropy and Bayesian methods*, Kluwer Academic Publishers, Boston.
- Heston S.L. (1993), *A Closed-Form Solution for Options with Stochastic Volatility with Applications to Bond and Currency Options*, "Review of Financial Studies", Vol. 6, pp. 327–343.
- Hull J., White A. (1987) *The pricing of options on assets with stochastic volatility*, "Journal of Finance", Vol. 42, pp. 281–300.
- Jackwerth J.C. (1999), *Option-implied risk-neutral distributions and implied binomial trees: A literature review*, "Journal of Derivatives", Vol. 7, No. 4, pp. 66–82.

- Jackwerth J.C. (2000), *Recovering Risk Aversion from Option Prices and Realized Returns*, "Review of Financial Studies", Vol. 13, pp. 433–451.
- Jarrow R., Rudd A. (1982), *Approximate valuation for arbitrary stochastic processes*, "Journal of Financial Economics", Vol. 10, pp. 347–369.
- Jondeau E., Rockinger M. (2000), *Reading the smile: the message conveyed by methods which infer risk neutral densities*, "Journal of International Money and Finance", Vol. 19, pp. 885–915.
- Jorion P. (1995), *Predicting Volatility in the Foreign Exchange Market*, "Journal of Finance", Vol. 50, No. 2, pp. 507–528.
- Leahy M.P., Thomas C.P.P. (1996), *The Sovereignty Option: The Quebec Referendum and Market Views on the Canadian Dollar*, "International Finance Discussion Paper", No. 555, Board of Governors of the Federal Reserve System.
- Madan D.B., Milne F. (1994), *Contingent claims valued and hedged by pricing and investing in a basis*, "Mathematical Finance", Vol. 4, pp. 223–245.
- Malz A.M. (1996), *Using option prices to estimate realignment probabilities in the European Monetary System: the case of sterling-mark*, "Journal of International Money and Finance", Vol. 15, October, pp. 717–748.
- Malz A.M. (1997), *Estimating the probability distribution of the future exchange rate from options prices*, "Journal of Derivatives", Winter, pp. 18–36.
- McDonald J.B. (1996), *Probability distributions for financial models*, w: G.S. Madalla, C.R. Rao (red.), *Statistical methods in finance*, Elsevier, Amsterdam.
- McDonald J.B., Bookstaber R.M. (1991), *Option pricing for generalized distributions*, "Communications in Statistics – Theory and Method", Vol. 20, pp. 4053–4068.
- Mc Manus D.J. (1999), *The information content of interest rate futures options*, "Working Paper", No. 99–15, Bank of Canada, Ottawa.
- Melick W.R. (1999), *Results of the estimation of implied PDFs from a common dataset*, in: BIS, *Estimating and Interpreting Probability Density Functions*, Proceedings of the workshop held at BIS on 14 June 1999, Basle.
- Melick W.R., Thomas, C.P. (1997), *Recovering an assets implied PDF from option prices: An application to crude oil during the Gulf crisis*, "Journal of Financial and Quantitative Analysis", Vol. 32, pp. 91–115.
- Merton R.C. (1976) *Option pricing when underlying returns are discontinuous*, "Journal of Financial Economics", Vol. 3, pp. 125–144.
- Micu M. (2005), *Extracting expectations from currency option prices: a comparison of methods*, paper presented at 11th International Conference "Computing in Economics and Finance 2005", Society for Computational Economics, 23–25 June, Washington, D.C., <http://ideas.repec.org/p/sce/scecf5/226.html>
- Mizrach B. (1996), *Did option prices predict the ERM crisis?*, "Working Paper", No. 96–10, Rutgers University, Department of Economics, New Brunswick.
- Navatte P., Villa C. (2000), *The information content of implied volatility, skewness and kurtosis: Empirical evidence from long-term CAC40 options*, "European Financial Management", Vol. 6, pp. 41–56.
- Rebonato R. (1999), *Volatility and correlations in the pricing of equity, FX and interest-rate options*, John Wiley, New York.
- Rosenberg J.V., Engle R.F. (2002), *Empirical Pricing Kernels*, "Journal of Financial Economics", Vol. 64, No. 3, pp. 341–372.
- Ross S. (1976), *Options and efficiency*, "Quarterly Journal of Economics", No. 90, pp. 75–89.
- Rubinstein M. (1994), *Implied binomial trees*, "Journal of Finance", Vol. 49, July, pp. 771–818.
- Samuelson P.A. (1965), *Rational Theory of Warrant Pricing*, "Industrial Management Review", Vol. 6, No. 2, pp. 13–31.
- Savickas R. (2001), *A simple option-pricing formula*, "Working paper", George Washington University, Department of Finance, Washington, D.C.
- Sherrick B.J., Garcia P.P., Tirupattur V. (1996), *Recovering Probabilistic Information from Option Markets: Tests of Distributional Assumptions*, "Journal of Futures Markets", Vol. 16, No. 5, pp. 545–560.
- Shimko D.C. (1993), *Bounds of probability*, "Risk", Vol. 6, April, pp. 33–37.
- Söderlind P.P., Svensson L.E.O. (1997), *New Techniques to Extract Market Expectations from Financial Instruments*, "Journal of Monetary Economics", Vol. 40, No. 2, pp. 383–429.
- Stephens M.A. (1974), *EDF Statistics for Goodness of Fit and Some Comparisons*, "Journal of the American Statistical Association", No. 69, pp. 730–737.
- Stopczyński A., Węgleńska N. (1999), *O wpływie świąt na cenę opcji*, "Rynek Terminowy", No. 3, pp. 92–95.
- Syrdal S.A. (2002), *A Study of Implied Risk-Neutral Density Functions in the Norwegian Option Market*, "Working Paper", No. 13, Norges Bank, Oslo.
- Weinberg S.A. (2001), *Interpreting the Volatility Smile: An Examination of the Information Content of Options Prices*, "International Finance Discussion Paper", No. 706, Board of Governors of the Federal Reserve System, Washington, D.C.

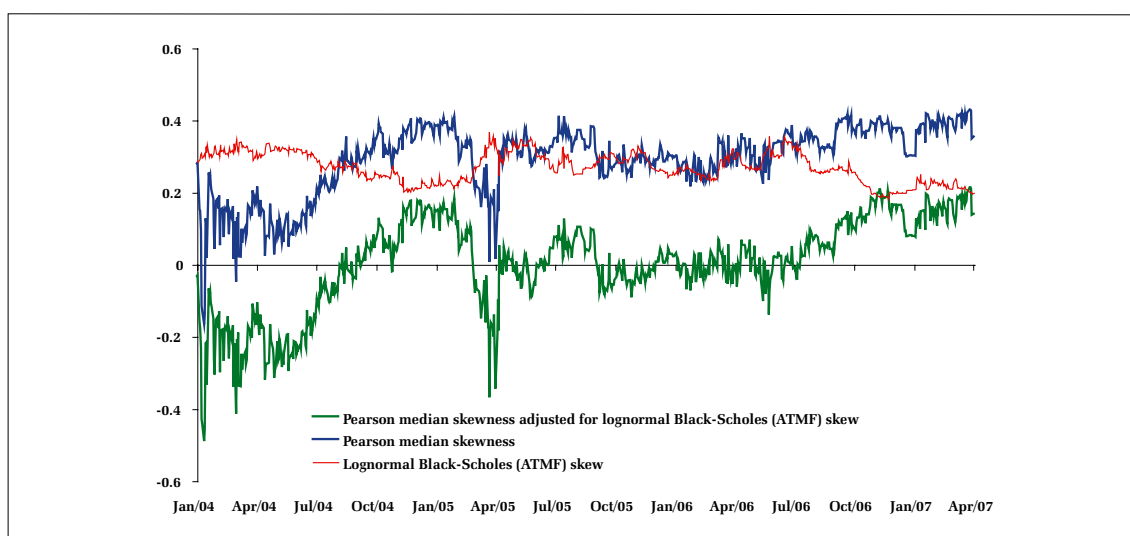
Appendix

Table A1. *Correlation of volatility measures (daily data)*

	PDF	ATMF	Historical
PDF	1	0.979	0.663
ATMF		1	0.660
Historical			1

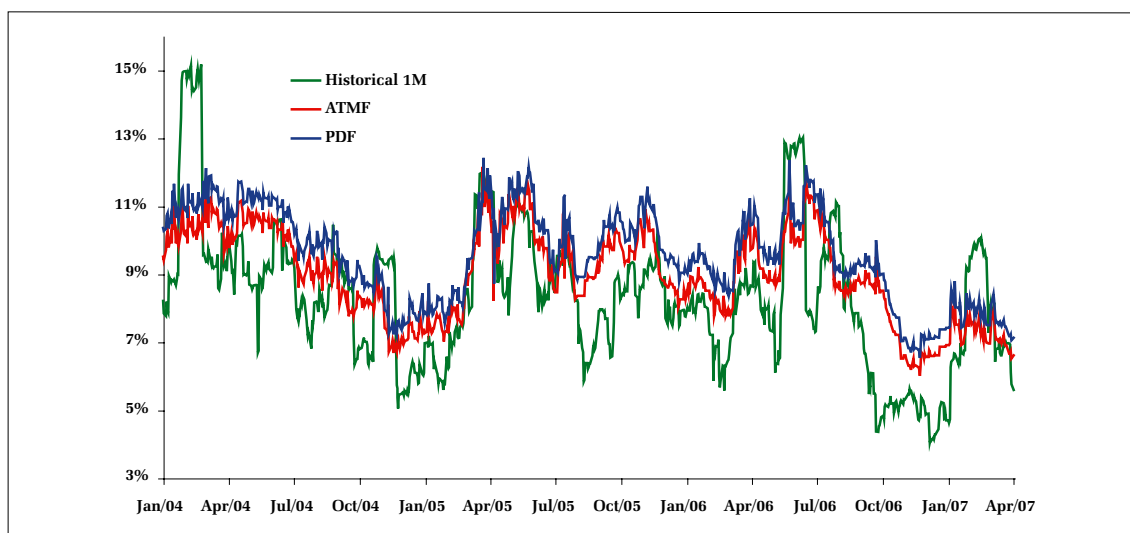
Source: Own calculations.

Figure A1. *Skewness measures*



Source: Own calculations.

Figure A2. *Volatility measures*



Source: Own calculations.

Obligacja jako narzędzie sekurytyzacji ryzyka śmiertelności

Bond as a Tool of Mortality Risk Securitization

*Aleksandra Matek**

pierwsza wersja: 17 stycznia 2008 r., ostateczna wersja: 13 marca 2008 r., akceptacja: 18 marca 2008 r.

Streszczenie

Ryzyko śmiertelności wiąże się zarówno ze śmiertelnością wyższą, jak i niższą niż oczekiwana.

Pierwszą kwestię należy wiązać przede wszystkim z terroryzmem, skutkami katastrof naturalnych i nieznanymi dotąd chorobami. Z kolei ryzyko wystąpienia śmiertelności niższej niż oczekiwana dotyczy produktów zabezpieczenia na starość, przy czym najważniejsze jest tzw. zagregowane ryzyko długowieczności – dotyczące ogółu populacji i świadczeń wypłacanych dożywotnio.

Idea transferu ryzyka ubezpieczeniowego na rynek kapitałowy narodziła się w branży majątkowej w odpowiedzi na kumulację szkód katastroficznych w latach 90. XX w. Nowe klasy ryzyka, weryfikowanie wymogów kapitałowych oraz popularność fuzji i przejęć spowodowały, że obligacje znajdują coraz większe zastosowanie w branży życiowej. Artykuł analizuje możliwości zastosowania obligacji do zarządzania ryzykiem śmiertelności.

Słowa kluczowe: sekurytyzacja, ryzyko śmiertelności, zarządzanie ryzykiem, obligacja katastrofowa, reasekuracja

Abstract

Mortality risk concerns both higher and lower than expected mortality developments.

The first aspect is linked to terrorism attacks, natural catastrophes and new classes of diseases, while the second one is focused primarily on retirement products. However, the crucial question of the second aspect is the aggregate longevity risk, concerning the overall population and annuities providers.

The idea of the insurance risk transfer has been developed by the P&C insurers and reinsurers, as an answer to catastrophic loss accumulation in the last decade of the 20th century. New classes of risk, capital requirement verification and M&A growth have created the possibility of using bonds in life insurance.

The article analyses the opportunities and possibilities of the usage of bonds as a tool of mortality risk securitization and risk management.

Keywords: securitization, mortality risk, risk management, catastrophe bond, reinsurance

JEL: G22, G23, G24

* Uniwersytet Ekonomiczny w Krakowie, Studium Doktoranckie, e-mail: aleksandra_malek@yahoo.co.uk

1. Wstęp

Ryzyko śmiertelności, rozumiane jako jej kształtowanie się w sposób odmienny aniżeli przewidywany, jest podstawową kwestią dla branży ubezpieczeń życiowych i musi być analizowane według dwóch scenariuszy. Pierwszy z nich wiąże się z wystąpieniem śmiertelności wyższej niż oczekiwana. W tym przypadku analizuje się ryzyko ekstremalnej śmiertelności. Drugi scenariusz dotyczy śmiertelności niższej niż oczekiwana i wtedy mówimy o ryzyku długowieczności.

Stopy śmiertelności mają tendencję do obniżania się, co jest pochodną poprawy jakości życia i postępu medycyny, a zatem pierwszy scenariusz należy wiązać przede wszystkim ze zdarzeniami o charakterze katastroficznym. Mają one szczególne znaczenie w dobie światowego terroryzmu, rozprzestrzeniania się nieznanych dotąd chorób, mogących skutkować epidemiami czy pandemiami, a także notowania niespotykanych dotąd skutków katastrof naturalnych. Wprawdzie ludzkość (a co za tym idzie – branża ubezpieczeń) zawsze zmagła się z chorobami oraz niszczącymi siłami przyrody, jednak dzisiaj zdarzenia te stanowią nową klasę ryzyka. Przede wszystkim katastrofy naturalne powodują ofiary śmiertelne głównie w społeczeństwach ubogich. W krajach rozwiniętych dzięki wysokim nakładom na badania naukowe i rozwój metod wczesnej identyfikacji kataklizmów szanse na ucieczkę są nieporównywalnie większe, a nagromadzenie szkód notowane jest zazwyczaj w branży majątkowej. Zagrożenie epidemiami, mimo postępu medycyny, poprawy higieny oraz akcji prewencyjnych, nadal jest aktualne. Niebezpieczeństwo tkwi z jednej strony w możliwości szybkiego i łatwego przemieszczania się ludzi (czego dowiódł SARS), a z drugiej – w mutacjach szczepów chorób (przykład ptasiej grypy). Z kolei ryzyko terroryzmu we współczesnej – globalnej – postaci powinno być postrzegane znacznie szerzej. Zagrożone są bowiem już nie tylko obszary konfliktów religijnych czy narodowościowych. Analiza musi także uwzględniać wykorzystanie broni biologicznej czy chemicznej.

Z kolei ryzyko wystąpienia śmiertelności niższej niż oczekiwana dotyczy przede wszystkim produktów zabezpieczenia na starość oraz ubezpieczeń zdrowotnych. Dla branży asekuracyjnej istotne jest tzw. zagregowane ryzyko długowieczności – dotyczące ogółu populacji i dotyczące świadczeń wypłacanych dożywotnio. Indywidualne ryzyko dotyczy natomiast pojedynczego ubezpieczonego i wiąże się z posiadaniem produktów terminowych (MacMinn et al. 2006, s. 551). Na skutek nabycia ochrony gwarantującej świadczenie dożywotnie, a także objęcia państwowym programem z wypłatami gwarantowanymi do chwili śmierci można wyeliminować ryzyko indywidualne poprzez scedowanie go na podmiot spełniający roszczenie. W związku z wydłużaniem się przeciętnego trwania życia i spadkiem przy-

rostu naturalnego ryzyko na poziomie zagregowanym staje się wyzwaniem właśnie dla tych podmiotów.

Obligacje umożliwiające transfer ryzyka ubezpieczeniowego na rynek kapitałowy były odpowiedzią na kumulację szkód w branży majątkowej wskutek megakatastrof lat 90. XX w. i wynikającą stąd konieczność poszukiwania nowych rozwiązań, dostarczających dodatkowej pojemności akceptacyjnej dla ryzyka i uniezależnionych od cykli reasekuracyjnych. Pojawianie się nowych klas ryzyka, weryfikowanie wymogów kapitałowych i popularność fuzji i przejęć spowodowały, że konstrukcja obligacji pierwotnie typowo katastrofowej¹ znajduje coraz szersze zastosowanie, zwłaszcza w ubezpieczeniach na życie.

Celem niniejszej pracy jest przedstawienie możliwości zaadaptowania obligacji do transferu ryzyka śmiertelności. Ze względu na rozległość tematu zdecydowano się jedynie skrótowo przedstawić zastosowania sekurytyzacji w ubezpieczeniach na życie, bez omawiania jej wad, zalet czy perspektyw rozwoju. Szerzej omówiono rolę sekurytyzacji w transferze typowych rodzajów ryzyka techniczno-ubezpieczeniowego, bazujących na instrumencie katastrofowym. Jako uzupełnienie teoretycznych rozważań przedstawione zostały konkretne programy sekurytyzujące ryzyko ekstremalnej śmiertelności oraz ryzyko długowieczności. Szczególną uwagę poświęcono strukturze indeksów jako bazy uruchamiania pokryć z instrumentów oraz konstrukcji wypłat.

2. Sekurytyzacja w branży ubezpieczeń

2.1. Zarys sekurytyzacji w branży życiowej

Zanim zostanie przedstawiona idea obligacji bazujących na ryzyku śmiertelności, konieczne jest dokonanie wprowadzenia do zagadnienia sekurytyzacji w branży życiowej. Mimo wielokierunkowości stosowanych rozwiązań można wyodrębnić jej trzy bazowe aspekty².

Programy VIF (*value in force*), znane też jako EV (*embedded-value*), umożliwiają monetyzację aktywów tzw. nieuchwytnych (Swiss Reinsurance Group 2006,

¹ Wykształciła się nawet terminologia związana z tego typu instrumentem – *Cat Bond*, *Catastrophic Bond*, *Catastrophe Bond*, *Act of God Bond*.

² Oprócz trzech głównych rodzajów sekurytyzacji w branży życiowej, którym poświęcono jest niniejszy podrozdział, A. Cowley i J. Cummins (2005, s. 208) wskazują dwa dodatkowe – sekurytyzację *viatical and life settlement* oraz sekurytyzację tzw. czystych aktywów. Drugi z nich jest sekurytyzacją typową dla przedsiębiorstw, zakładów zdrowotnych i banków w celu (re)finansowania ich działalności. Z powodu rozległości tematu, a także wielości typów sekurytyzacji w branży asekuracyjnej to zagadnienie nie jest poruszane w niniejszej pracy. *Viatical and life settlement* wiąże się z wykupywaniem polis od ich posiadaczy, zazwyczaj przez brokera albo przedsiębiorców, co umożliwia otrzymanie środków przed wykupem polisy. Następnie te „wykupione” polisy mogą być poddane sekurytyzacji i oferowane inwestorom. Ze względu na szczególne cele takiego zabiegu i fakt, że sponsorem takiej sekurytyzacji jest nie zakład ubezpieczeń sprzedający polisy, lecz pośrednik (broker lub przedsiębiorca), można uznać, iż operacja nie jest narzędziem zarządzania ryzykiem ubezpieczyciela i finansowania jego działalności. Nie należy więc do głównych trzech przejawów sekurytyzacji, o których mowa powyżej, i są nieistotne w świetle zagadnień poruszanych w niniejszej pracy. Por. Stone, Zissu (2006).

s. 8), takich jak rozkładane w czasie, odroczone koszty akwizycji czy wartość bieżąca przyszłych zysków³. W klasycznym wariantcie – bez sekurytyzacji – koszty są amortyzowane przychodami ze składek wnoszonych w okresie ubezpieczenia, co wiąże się z płynnością. Strategia ta stale jest także obciążona ryzykiem wcześniejszej rezygnacji przez ubezpieczonego ze świadczonej ochrony, a osiągnięte zyski zależą dodatkowo od śmiertelności i poziomu stóp procentowych. Mimo że ten typ sekurytyzacji uważa się raczej za narzędzie finansowania działalności niż hedging, należy podkreślić, że dotyczy on także transferu ryzyka, a więc ingeruje w zarządzanie nim. Oczywiście część ryzyka transferowanego w transakcjach tego typu to ryzyko techniczno-ubezpieczeniowe, w szczególności wyższej niż oczekiwana śmiertelności lub niższej niż oczekiwana, gdy sekurytyzacja dotyczy produktów zabezpieczenia na starość (Swiss Reinsurance Group 2006, s. 8–9). W tym momencie należy poczynić kilka uwag. Przede wszystkim warto zauważyć, że omawiany transfer jest niejako skutkiem ubocznym specyfiki ryzyka śmiertelności, będącego przecież głównym ryzykiem techniczno-ubezpieczeniowym w branży *life*, a nie podstawowym celem przeprowadzanej transakcji. Poza tym praktyka sekurytyzacji w asekuracji spowodowała powstanie narzędzia typowego alternatywnego transferu ryzyka techniczno-ubezpieczeniowego – mowa tu o trzecim aspekcie sekurytyzacji, omawianym poniżej. Rozwiązania te wykorzystywane były pierwotnie w branży majątkowej. Na inwestorów (czy pierwotnie na spółkę specjalnego przeznaczenia) cedowany jest pakiet różnych rodzajów ryzyka, mających wpływ na uzyskiwany wynik, podczas gdy alternatywny transfer obejmuje jedynie ryzyko techniczno-ubezpieczeniowe (wzmózonej śmiertelności – skutkującej kumulacją świadczeń z tytułu ubezpieczeń na życie – lub długowieczności związanej z zabezpieczeniami emerytalno-rentowymi).

Jeden z pierwszych programów VIF wprowadził Hannover Re, zyskując w 1998 r. dzięki transakcji L1 o wolumenie 51 mln EUR możliwość sfinansowania działalności związanej z nowymi rodzajami ubezpieczeń w Zachodniej Europie. W latach 1999–2002 reasekurator kontynuował strategię, włączając do programu kolejno ubezpieczenia zdrowotne i wypadkowe istniejących (128 mln EUR w ramach L2) i nowych (50 mln EUR w ramach L3) portfeli, a także powiązane z funduszami kapitałowymi: L4 o wolumenie 200 mln EUR i L5 – 300 mln EUR (Bütow 2001, s. 23)⁴.

Druga strategia wiąże się z koniecznością spełniania wymogów kapitałowych. Transakcje określane mianem sekurytyzacji XXX i AXXX są stosowane na rynku amerykańskim na skutek wprowadzenia regulacji utrzymujących konserwatywne założenia wyceny, niemające

współcześnie uzasadnienia ekonomicznego⁵. Zabieg sekurytyzacji skutkuje emisją obligacji o wartości równej „zbędnym” rezerwom, stanowiącym różnicę między wartościami ustawowymi a rzeczywistymi, uzasadnionymi ekonomicznie. W wyniku tego zmniejsza się „obciążenie” rezerwami i jednocześnie są one zabezpieczone środkami wniesionymi przez inwestorów. Środki te, należne z tytułu nabycia instrumentów, utrzymuje spółka specjalnego przeznaczenia (SPV) (Swiss Reinsurance Group 2006, s. 11⁶). W przypadku konieczności wykorzystania „pierwotnych” rezerw, zastąpionych sekurytyzacją, a więc np. gdy śmiertelność jest wyższa niż oczekiwana, uruchomione zostaną środki w celu spełnienia tych „nadwyżkowych” roszczeń. Do SPV trafiają również opłaty uzyskiwane od sponsora emisji, które następnie są swapowane⁷, a uzyskiwane w ten sposób oprocentowanie, zależne od stóp międzybankowych, wypłacane jest inwestorom⁸. W transakcjach tego typu może także występować gwarant. W zależności od założeń programu gwarantowane są poszczególne transze lub wszystkie instrumenty, co zapewnia zwrot wniesionej wartości nominalnej w części bądź w całości. Sekurytyzacja rezerw znalazła zastosowanie m.in. w First Colony Life Insurance Company; w czterech transakcjach w latach 2003–2004 dostarczyła łącznie 1,15 mln USD (Cowley, Cummins 2005, s. 218).

Trzecim typem sekurytyzacji istotnym w niniejszym opracowaniu jest alternatywny transfer ryzyka. Opiera się on na rozwiązaniach wypracowanych w branży majątkowej i polega na uzyskaniu na rynku kapitałowym dodatkowego, z reguły korzystniejszego niż na tradycyjnym rynku (re)asekuracyjnym, pokrycia ryzyka ubezpieczeniowego. Formalnie zabieg ten dotyczy zatem pasywów, a ściślej – przyszłych zobowiązań z tytułu sprzedanych polis, warunkowanych przebiegiem szkodowości czy zajściem zdarzenia ubezpieczeniowego. Rozwój technik sekurytyzacyjnych w ubezpieczeniach jest wynikiem dającej się zaobserwować na rynku finansowym ogólnej tendencji do wykorzystywania innowacji finansowych. Sekurytyzacja ubezpieczeniowa pozostaje jednak specyficznym instrumentem pomimo istnienia niezaprzeczalnej analogii między procesami występującymi w branży (re)asekuracyjnej i bankowej – bezpośrednim bodźcem do rozwoju sekurytyzacji bankowej

³ Ten typ wykorzystywany jest także w operacjach demutualizacji i przejęć – por. Cowley, Cummins (2005, s. 212–218); PricewaterhouseCoopers (2006, s. 2).

⁴ W kwestii L5 por. Cowley, Cummins (2005, s. 211).

⁵ Na temat szczegółów dotyczących Regulacji XXX i AXXX por. Swiss Reinsurance Group (2006, s. 11).

⁶ W opracowaniu omówiono także typową strukturę sekurytyzacji EV.

⁷ „Swapowanie” jest terminem wywodzącym się od słowa „swap” – instrumentu pochodnego, którego istotą jest wymiana przepływów pieniężnych pomiędzy stronami umowy (kontraktu) i równoważne jest określeniu „zawrzeć swap”, ewentualnie „wstąpić w swap”. W analizowanych przypadkach, o ile nie zaznaczono inaczej, mowa jest o swapie na stopę procentową. Wymiana płatności ma tu na celu zagwarantowanie inwestorom kuponu bazującego na zmieniającym oprocentowaniu.

⁸ Por. Cowley, Cummins (2005, rys. 11).

było zmniejszenie pojemności kredytowej⁹ na skutek ogólnoswiatowego kryzysu na rynkach dłużnych w latach 80. ubiegłego stulecia (Swiss Reinsurance Group 1996, s. 6). Programy w branży niezyciowej mają na celu głównie zabezpieczenie przed skutkami katastrof – zarówno naturalnych, jak i spowodowanych przez człowieka (tzw. *man-made disasters*). Z kolei zakłady ubezpieczeń na życie, czerpiąc z tych doświadczeń, wypracowały nowe rozwiązania, transferujące na rynek kapitałowy ryzyko śmiertelności.

2.2. Istota sekurytyzacji transferującej ryzyko

Instrumenty sekurytyzacji transferujące ryzyko ubezpieczeniowe określane są terminem *insurance/risk-linked securities* (ILS/RLS). Lwia część rynku ILS przypada na emisje obligacji¹⁰.

Można wyróżnić dwie główne strategie sekurytyzacji (Hartung 2004, s. 401; Albrecht, Schradin 1998, s. 579): bezpośrednią (pierwotną) i pośrednią (wtórną). W pierwszym przypadku podmiotem transferującym ryzyko jest cedent (zakład ubezpieczeń)¹¹. Z reguły jest on jedynym emitentem papierów wartościowych, a środki uzyskane z emisji lokuje w trakcie trwania transakcji w inwestycje wolne od ryzyka. Często pośredniczy w tym bank, doradczający przy emisji papierów i ich sprzedaży. Zastosowanie sekurytyzacji pierwotnej, oprócz niewątpliwej zalety polegającej na dostarczeniu

kapitału, ma również wiele wad, przede wszystkim w kwestii oddziaływania na bilans. Niezaprzeczalną zaletą sekurytyzacji bezpośredniej jest możliwość nawiązania ściślejszych kontaktów z innymi podmiotami z sektora finansowego (Bütow 2001, s. 6)¹².

Znacznie popularniejsza sekurytyzacja wtórna (pośrednia) polega z kolei na założeniu niezależnej spółki celowej: *special purpose vehicle* – SPV lub *special purpose reinsurance vehicle* – SPR (Ronka-Chmielowiec 2004, s. 215; Punter 2000, s. 9), mającej za zadanie przejęcie ryzyka w ramach konwencjonalnych kontraktów reasekuracyjnych i jednocześnie refinansowanie się poprzez emisję papierów wartościowych (por. schemat 1).

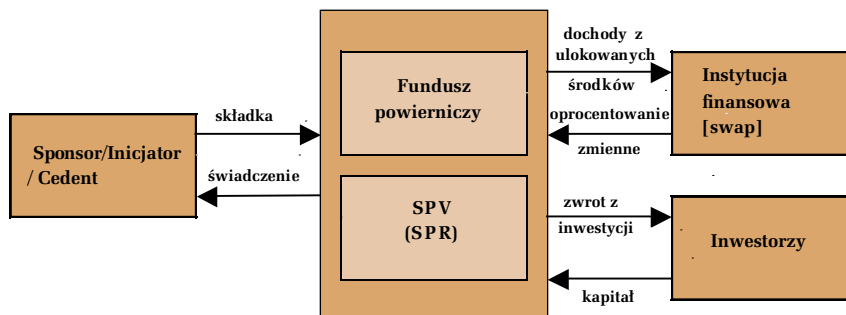
¹² Podział sekurytyzacji na „pośrednią” („wtórną”) i „bezpośrednią” („pierwotną”) występuje przede wszystkim w literaturze niemieckojęzycznej. Konstrukcja sekurytyzacji pierwotnej bez podmiotu SPV, do którego – zgodnie z definicją sekurytyzacji – powinno być cedowane ryzyko, może być podana w wątpliwość. Może z niej bowiem wynikać, że sekurytyzacją „zwykłą” zostanie nazwana emisja obligacji. Wobec ewentualnych zarzutów nadużywania określenia „sekurytyzacja” czy bezpodstawnego używania go w przypadku programów nieuwzględniających SPV autorka pragnie dodać, że w literaturze przedmiotu nie ma obecnie rozgraniczenia, co powinno podlegać w sferze alternatywnego transferu ryzyka na rynek kapitałowy sekurytyzacji zgodnie z jej definicją i jak należy się odnieść do wzorców OECD czy NUK. W dalszej części artykułu autorka wskazuje, że znacznie popularniejsza jest sekurytyzacja pośrednia, i na niej opiera dalszą analizę obligacji na ryzyko wzmożonej śmiertelności. Koncepcja sekurytyzacji bezpośredniej ryzyka śmiertelności jest jednak „wygodna” i wydaje się nie do obalenia. Świadczy o tym np. definiowanie obligacji na długowieczność właśnie jako „sekurytyzacji ryzyka długowieczności” (por. obligacja wyemitowana przez Europejski Bank Inwestycyjny, podrozdział 4.). Zdaniem autorki nazywanie sekurytyzacją zdecydowanej większości instrumentów alternatywnego transferu ryzyka (także instrumentów warunkowego kapitału, według niektórych autorów) nie do końca jest zgodne nie tylko z praktyką bankowej sekurytyzacji, ale także z praktyką tej asekuracyjnej. Istnieje zresztą problem z nazwaniem instrumentów alternatywnego transferu ryzyka (ART). Część autorów wyróżnia tzw. ILS (*insurance-linked securities*) i ubezpieczeniowe instrumenty pochodne (*insurance derivatives*), głównie w Niemczech. Część z kolei wszystkie te instrumenty klasyfikuje jako ILS, zarówno pochodne, jak i hybrydy. Autorka rozważała stworzenie na potrzeby tego opracowania dwóch klas: sekurytyzacji *sensu largo* i sekurytyzacji *sensu stricto*, jednak to nie rozwiązuje problemu. Z kolei wyodrębnienie klasy „sekurytyzacji na potrzeby asekuracji” mogłoby ten problem rozwiązać, jednak autorka chciałaby uniknąć rozgraniczeń „sekurytyzacja na potrzeby asekuracji” i „sekurytyzacja [inna? zwykła? pozaasekuracyjna?]”. Podważałoby to bowiem konwergencję rynków: asekuracyjnego i kapitałowego, która jest przecież istotą i racją bytu dla produktów ART.

⁹ Pojemność kredytowa rozumiana jest jako górna granica akcji kredytowej banków (banku). Granica ta jest pochodną między innymi regulacji ostrożnościowych organu nadzoru i wewnętrznej polityki poszczególnych banków (w przypadku rynku polskiego coraz częściej grupy kapitałowej, banku będącego głównym udziałowcem etc.) w zakresie fundingu, utrzymywania kapitałów w stosunku do udzielanych pożyczek lub kredytów, koncentracji w stosunku do pojedynczego podmiotu bądź grupy podmiotów.

¹⁰ Więcej na ten temat: Swiss Re (2006, s. 24–29).

¹¹ Oczywiście w roli podmiotu transferującego może występować także reasekurator. Dla porządku jednak należy dodać, że w takim przypadku jest on już nie cedentem, a retrocedentem.

Schemat 1. Sekurytyzacja pośrednia



Na skutek unormowań prawnych niezwykle popularne staje się plasowanie SPV w ramach podatkowych, a sama spółka ma cechy towarzystwa zależnego, tzw. *captive* (Strube 2001, s. 53; Ronka-Chmielowiec 2004, s. 215)¹³.

Występowanie w roli SPV towarzystwa reasekuracji (czy brokera reasekuracyjnego – por. Wagner 1999, s. 12) wynika z konieczności zdefiniowania jednoznacznego prawnie stosunku pomiędzy cedentem a SPV oraz zaakceptowania transferu przez nadzór (Ronka-Chmielowiec 2004, s. 216). Specyfika obowiązków SPV jest pochodną niezwykle dokładnych restrykcji odnoszących się do wypłacalności podmiotu. Może ona inwestować jedynie w bezpieczne papiery i nie może prowadzić innej działalności niż zarządzanie wniesionymi środkami. Do utworzenia spółki niezbędne okazują się z reguły banki inwestycyjne, organizujące emisję papierów i rozwijające struktury spółki, a także profesjonalni pośrednicy finansowi – przede wszystkim mający doświadczenie w finansowaniu typu *venture*. Zamiast cedowania wszystkich funkcji na pośrednika można powołać większą liczbę podmiotów o specjalnych kompetencjach – np. reasekuratora bądź brokera reasekuracyjnego oraz bank albo maklera (Strube 2001, s. 54; Wagner 1997, s. 12).

Dotychczasowe sekurytyzacje ciągle są zabezpieczeniem przed skutkami katastrof¹⁴, tymczasem zakłady ubezpieczeń na życie poszukują nowych możliwości zabezpieczania się przed ryzykiem ekstremalnej długowieczności i ryzykiem wzmożonej śmiertelności. Tym zagadnieniom zostaną poświęcone dalsze części pracy.

3. Obligacje bazujące na ryzyku wzmożonej śmiertelności

3.1. Programy sekurytyzacji

Dotychczas przeprowadzono pięć udanych programów sekurytyzacji ryzyka wzmożonej śmiertelności. Obligacje te wykorzystują rozwiązania wypracowane w branży majątkowej w odniesieniu do ryzyka katastrof¹⁵. Opieranie się na sprawdzonych wzorcach i jednocześnie stosowanie dywersyfikacji kluczowej dla ryzyka śmiertelności przyczyniły się do sukcesu alternatywnego transferu ryzyka właściwego dla branży życiowej.

Do końca 2007 r. przeprowadzono pięć programów

sekurytyzacji ryzyka wzmożonej śmiertelności. W czterech z nich w roli sponsora¹⁶ występowali reasekuratorzy: Swiss Re (programy Vita Capital I, II i III) i Scottish Re (program Tartan Capital). Jedynym ubezpieczycielem sekurytyzującym portfele ryzyka śmiertelności była Axa (program Osiris Capital). Podstawowe parametry programów przedstawia tabela 1. W celu wykorzystania możliwości stworzenia kompletnego instrumentu zarządzania ryzykiem obligacje emitowane są w transzach (zwyczajowo zwanych klasami, zazwyczaj oznaczanych literami alfabetu łacińskiego i cyframi arabskimi, rzadziej rzymskimi), zróżnicowanych pod względem ryzyka dla inwestorów. Zróżnicowanie to osiągnięte jest przede wszystkim w wyniku ustanowienia różnych poziomów uruchamiania pokrycia dla cedenta¹⁷ (por. kolumna 3 w tabeli 1), rozumianego jako kompensata z tytułu nagromadzenia roszczeń będących rezultatem wystąpienia ekstremalnej śmiertelności. W przypadku analizowanych programów mechanizm uruchamiania pokrycia (MUP, ang. *trigger*) opiera się na wartościach zdefiniowanego wcześniej indeksu populacji, zdyspersyfikowanego pod względem płci, wieku i geograficznie (wyjątkiem od tej reguły jest Tartan Capital). Różne poziomy uruchamiania pokryć w poszczególnych transzach umożliwiają stworzenie kompletnego programu – limit pokrycia danej transzy jest automatycznie dolną granicą uruchomienia pokrycia z transzy kolejnej pod względem wymagalności. Dzięki temu eliminuje się luki w zabezpieczeniu – wyczerpanie środków inwestorów w transzy najbardziej ryzykownej (pierwszej pod względem wymagalności) skutkuje automatycznym wykorzystaniem (w razie potrzeby) środków inwestorów mających instrumenty wyemitowane w transzach kolejnych pod względem wymagalności (mniej ryzykownych) transzach, aż do osiągnięcia limitu transzy, która jest ostatnią pod względem wymagalności. Pokryciem są środki inwestorów – w zależności od konstrukcji instrumentu z wniesionych przez nich wartości nominalnych bądź (i) należnych im kuponów. Oprocentowanie poszczególnych transz zależy od poziomów uruchamiania pokrycia, a także (w mniejszym stopniu) okresu zapadalności. Ta druga zależność jest szczególnie widoczna w przypadku programu Vita Capital III: transze A-IV i A-V różnią się jedynie okresem zapadalności, co powoduje nieznaczne różnice w oprocentowaniu. Rating obligacji ustalany jest odrębnie dla poszczególnych transz.

W roli sponsora pierwszego programu sekurytyzacji wystąpiło w grudniu 2003 r. Swiss Re. Dzięki wyemitowaniu obligacji z zagrożoną wartością nominalną za pośrednictwem spółki specjalnego przeznaczenia Vita Capital zyskało pokrycie w wysokości 400 mln USD (Swiss Re 2003). Sukces programu skłonił reasekuratora do przeprowadzenia w kwietniu 2005 r. kolejnej sekurytyzacji – w wysokości 362 mln USD. Podobnie jak

¹³ Coraz częściej spółki wykorzystują formę Protected Cell Company (PCC) – por. Bütow 2001, s. 9.

¹⁴ Wyjątkami są nieliczne sekurytyzacje zobowiązań szkód z ubezpieczeń motoryzacyjnych czy wypadków przemysłowych (por. Swiss Reinsurance Group 2006, s. 38–40).

¹⁵ Szczegóły dotyczące typowej transakcji sekurytyzacji katastrof (a więc sekurytyzacji wtórnej – por. Bütow 2001, s. 6–7) znajdują się m.in. w opracowaniu Swiss Reinsurance Group (2006, s. 8), na temat konstrukcji konkretnego programu por. np. Blake et al. (2006a, s. 8) oraz Cowley, Cummins (2005, s. 222). Struktura czasowa przepływów pieniężnych z tytułu obligacji zobrazowana jest m.in. w: Małek (2007b, s. 30).

¹⁶ Sponsor i inicjator używane są zamiennie w niniejszym opracowaniu.

¹⁷ Cedenta ryzyka poddawane sekurytyzacji, ale w rzeczywistości retrocedenta dla tego ryzyka, ze względu na wcześniejszy stosunek reasekuracji.

Tabela 1. Instrumenty wyemitowane w ramach programów sekurytyzacji ryzyka wzmożonej śmiertelności

Program	Transza	Moment uruchamiania pokrycia (ang. <i>trigger</i>) w %	Górna granica pokrycia (limit) w %	Oprocentowanie obligacji	Okres zapadalności	Rating (S&P/Moody's)
Vita Capital	400 mln USD	130	150	LIBOR 3M + 135 pb	4 lata	A+ / A3
Vita Capital II	B – 62mln USD	120	125	LIBOR 3M + 90 pb	5 lat	A- /Aa3
	C – 200 mln USD	115	120	LIBOR 3M + 140 pb		BBB+ / A2
	D – 100 mln USD	110	115	LIBOR 3M + 190 pb		BBB- / Baa2
Osiris Capital	B1 – 100 mln EUR	114	119	EURIBOR 3M + 20 pb	4 lata	AAA / Aaa
	B2 – 50 mln EUR	114	119	EURIBOR 3M + 120 pb		A- / A3
	C – 150 mln USD	110	114	LIBOR 3M + 285 pb		BBB / Baa2
	D – 100 mln USD	106	110	LIBOR 3M + 500 pb		BB+ / Ba1
Tartan Capital	A – 75 mln USD	115	120	LIBOR 3M + 19 pb	3 lata	AAA / Aaa
	B – 80 mln USD	110	115	LIBOR 3 M + 300 pb		BBB / Baa3
Vita Capital III	A-IV – 100 mln USD	125	145	LIBOR 3M + 21 pb	4 lata	AAA/ Aaa
	A-V – 100 mln USD			LIBOR 3M + 20 pb	5 lat	AAA/ Aaa
	A-VI – 55 mln EUR			EURIBOR 3M + 21 pb	4 lata	AAA/ Aaa
	A-VII – 100 mln EUR			EURIBOR 3M+ 80 pb	5 lat	AA-/Aa2
	B-I – 90 mln USD	120	125	LIBOR 3M + 110 pb	4 lata	A/A1
	B-II – 50 mln USD			LIBOR 3M + 112 pb	5 lat	A/A1
	B-III – 50 mln EUR			EURIBOR 3M + 110 pb	4 lata	A/A1
	B-V – 50 mln USD			LIBOR 3M + 21 pb	5 lat	AAA/ Aaa
	B-VI – 55 mln EUR			EURIBOR 3M + 22 pb	4 lata	AAA/ Aaa

Źródło: opracowanie własne na podstawie: Swiss Reinsurance Group (2006, s. 37); Axa (2006, s. 1); Bauer, Kramer (2008, s. 6).

w pierwszym przypadku wyemitowane instrumenty zakładały możliwość utraty wartości nominalnej (tzw. *principal at risk bond*), jednak struktura wypłat bazowała na opisanym powyżej mechanizmie różnicowania transz-klas (por. tabela 1). W styczniu 2007 r. Swiss Re uzupełniło uzyskane na rynku kapitałowym pokrycie o kolejne 705 mln USD, emitując instrumenty w siedmiu transzach (Swiss Re 2007).

Opisane powyżej programy stanowią w rzeczywistości alternatywę dla tradycyjnej retrocesji¹⁸. W ramach pierwszej (listopad 2006 r.) i na razie jedynej sekurytyzacji stanowiącej alternatywę dla reasekuracji zostały wyemitowane instrumenty transferujące ryzyko od Axy.

Strukturę programu przedstawia schemat 2. Analogicznie do klasycznego stosunku reasekuracji za przyjęcie ryzyka sponsor musi uiścić składki na rzecz spółki specjalnego przeznaczenia (SPV – Osiris Capital PLC), formalnie występującej w roli reasekuratora. W celu zapewnienia bezpieczeństwa przepływów pieniężnych „przechodzących” przez SPV wykorzystuje się rachunki powiernicze, na których deponowane są środki wno-

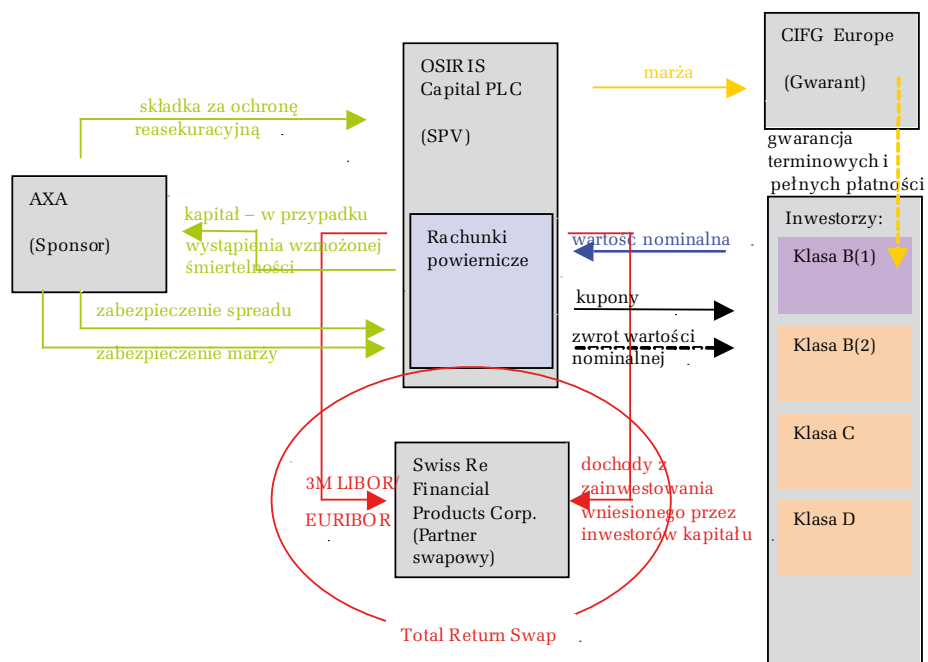
szone przez inwestorów z tytułu nabycia papierów, jak też dokonywane przez Axę płatności zabezpieczające *spread* ponad stopę rynku międzybankowego oraz marżę dla gwaranta. *Spread* jest głównym elementem różnicującym oprocentowanie transz. Zgodnie z praktyką różnicowania ryzyka, opisaną na początku podrozdziału, transze instrumentów, które są pierwsze pod względem wymagalności, charakteryzuje wyższy *spread*, przy czym maleje on wprost proporcjonalnie do redukcji ryzyka utraty środków, aż do wartości minimalnej w przypadku transz z gwarancją otrzymania kapitału (i odsetek)¹⁹. Program Osiris Capital obejmuje cztery klasy instrumentów o łącznej wartości nominalnej 442 mln USD (345 mln EUR)²⁰, zróżnicowane pod względem walut i ryzyka, zawierające opcję utraty wartości nominalnej. Wyjątek stanowią wyemitowane w ramach klasy B papiery o wartości 100 mln EUR, które zabezpiecza przed utratą CIFG Europe i z tego względu otrzymały rating AAA/Aaa. Taka konstrukcja skutkuje obniżeniem oprocentowania transzy – dochody inwestorów wy-

¹⁸ Retrocesja polega na cedowaniu ryzyka reasekuratora na inny podmiot. Pod względem technicznym jest więc ubezpieczeniem reasekuratora, a zatem kolejnym stadium wtórnego transferu ryzyka ubezpieczeniowego.

¹⁹ Mimo gwarancji zwrotu kapitału (i najczęściej także wypłaty kuponów) oprocentowanie transzy nie jest równe stopie referencyjnej – *spread* ponad LIBOR/EURIBOR odzwierciedla premię za płynność oraz (ciągle jeszcze) premię za nowość (tzw. *novelty premium*).

²⁰ Czasami podaje się kwoty w dolarach i w przeliczeniu na euro. Wynika to z tego, że część transzy emitowana jest w dolarach, a część w euro.

Schemat 2. Struktura programu Osiris Capital PLC



Źródło: opracowanie własne na podstawie: Standard&Poors, Osiris Capital PLC. Commentary Report, s. 5.

znaczone są jako 3M EURIBOR + 20 bp, podczas gdy oprocentowanie transzy, która jest druga pod względem bezpieczeństwa wnoszonych środków (również B, ale w wariantcie *principal at risk* (rating A-/A3)) wynosi 120 bp ponad 3-miesięczny EURIBOR. Najbardziej ryzykowna klasa D, w której zostały wyemitowane papiery o wolumenie 100 mln USD (które jako pierwsze mogą być utracone w przypadku wystąpienia zdarzenia uruchamiającego pokrycie) wypłaca kupony zależne od wysokości LIBOR 3M powiększonej o 500 bp. Klasa C, druga pod względem wymagalności pokrycia (rating BBB/Baa2), wypłaca 285 punktów bazowych ponad trzymiesięczny LIBOR (Axa 2006, s. 1).

3.2. Indeks jako istota programu sekurytyzacji

Dla każdego programu sekurytyzacji podstawową kwestią są warunki pokrycia i moment jego uruchomienia (MUP, ang. *trigger*). W wyniku wieloletnich doświadczeń w branży majątkowej wykształciły się różnorodne rozwiązania dotyczące MUP – od opartych na szkodach w portfelach sponsora (wyrażonych jako wartość wszystkich szkód bądź liczba roszczeń), fizycznych parametrach katastrofy (tzw. MUP parametryczny, opierający się przykładowo na ciśnieniu w oku cyklonu, sile trzęsienia ziemi, prędkości wiatru), poprzez indeksy branżowe, regionalne, parametryczne, po straty mo-

delowane. W przypadku programów sekurytyzacji ryzyka ekstremalnej śmiertelności podstawą uruchomienia pokrycia dla sponsorów jest wartość indeksu bazująca (mniej lub bardziej pośrednio) na śmiertelności portfeli, uzbrojona jednak w indeks skonstruowany specjalnie na potrzeby pojedynczego programu, dywersyfikowany terytorialnie, wiekowo i pod względem płci.

W dotychczasowych programach branży majątkowej MUP stanowi kwestię dyskusyjną, co wynika z różnicy interesów sponsora i inwestorów. Rozwiązania zapewniające przejrzystość wykorzystywanych danych (dzięki czemu łatwo zmniejsza się ryzyko hazardu moralnego sponsora), a więc niezależne od sytuacji sponsora – z reguły oparte na indeksach czy parametrach katastrofy – nie zawsze zapewniają dopasowanie zabezpieczenia do rzeczywistych strat w posiadanych portfelach i są źródłem istotnego ryzyka bazy²¹. Rozwiązanie oparte na wielowymiarowej dywersyfikacji indeksu, opracowane przez Swiss Re w programie Vita Capital i stosowane z sukcesem w kolejnych programach reasekuratora, a także w przypadku Osiris Capital, pozwala na uwzględnienie przejrzystości i jednocześnie zaspokojenie potrzeb „portfelowych” sponsora.

Zastosowanie dywersyfikacji terytorialnej, oprócz oczywistego dla inwestorów rozproszenia ryzyka, za-

²¹ Szerzej nt. zastosowań i właściwości MUP w branży majątkowej: Dubinski, Laster (2003); Wagner (1997, s. 18-19).

Tabela 2. Pandemie XX w. i spowodowany nimi wzrost śmiertelności w grupie osób poniżej 65. roku życia (w %)

Rok	Wzrost śmiertelności w grupie osób poniżej 65. roku życia
1918	ponad 90
1936–1937	około 60
1943–1944	około 30
1957–1958	36
1967–1968	4
1968–1969	około 40
od roku 1992	poniżej 10

Źródło: opracowanie własne na podstawie: Monday (2006, s. 24).

pewnia najwyższy obiektywizm w szacowaniu szkód dzięki wykorzystaniu danych z oficjalnych źródeł. Nie bez znaczenia jest także możliwość szybkiej weryfikacji, czy konieczne jest uruchomienie pokrycia²². Przykładowo, w przypadku Vita Capital indeks obejmujący terytorium USA, Wielkiej Brytanii, Francji, Szwajcarii i Włoch (wagi odpowiednio: 70%, 15%, 7,5%, 5%, 2,5%) opierał się na danych dostarczanych przez organy rządowe i urzędy statystyczne (Cowley, Cummins 2005, s. 222).

Nadanie wag ma na celu uwzględnienie potrzeb portfelowych sponsora. Sytuacja ta jest szczególnie widoczna w przypadku programu Osiris Capital, w którym konstrukcja indeksu nie odzwierciedla rzeczywistych proporcji populacji stanowiących jego podstawę, lecz skupia się raczej na geograficznym rozkładzie ryzyka Axy. Śmiertelność we Francji wpływa na indeks aż w 60%; poza tym wpływa na niego śmiertelność zanotowana w Japonii – w 25% – i w USA – w 15% (Axa

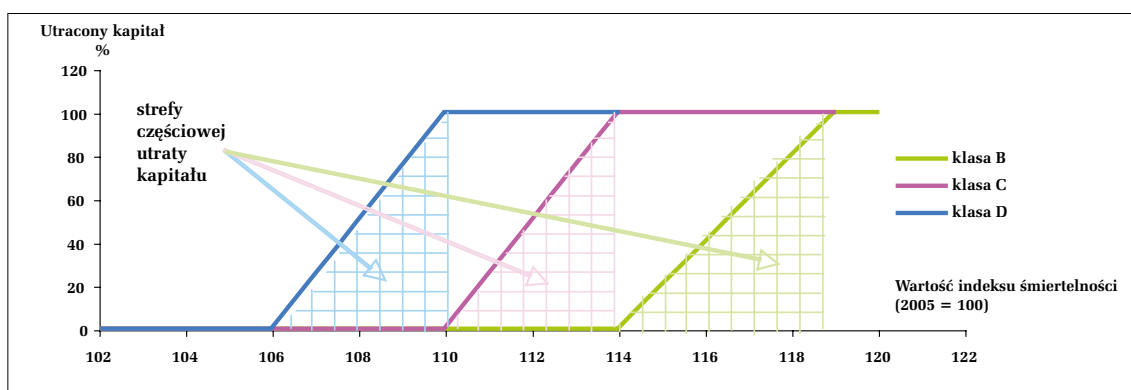
2006, s. 1). Sensowności i efektywności dywersyfikacji terytorialnej dowodzi przypadek Tartan Capital. Wypłata z obligacji sponsorowanych przez Scottish Re zależy od przebiegu śmiertelności jedynie na terenie USA – kraju szczególnie narażonego na ryzyko zamachów terrorystycznych. Emisja doszła do skutku, jednak w porównaniu z pozostałymi sekurytyzacjami zainteresowanie sponsorów było niewielkie. Podczas gdy w przypadku Vita Capital dochodziło do nadsubskrypcji, Scottish Re pozyskało jedynie 155 mln USD wobec oczekiwanych 200 mln USD (Reactions 2006).

Dywersyfikacja wiekowa jest pochodną specyfiki ryzyka ekstremalnej śmiertelności. Wysokie wagi przypisane przedziałom wiekowym 30–60 wiążą się głównie z ryzykiem aktów terrorystycznych, których obiektem są osoby aktywne zawodowe, relatywnie zamożne (i jednocześnie zazwyczaj mające polisę na życie). Te same osoby charakteryzują się jednocześnie mniejszym ryzykiem śmiertelności wskutek pandemii (por. tabela 2). Umożliwia to wewnętrzną neutralizację ryzyka.

Warto również wspomnieć, że indeks nie musi odnosić się do okresu 1 roku. Rozwiązanie zapoczątkowane przez Vita Capital II, oparte na śmiertelności w dwóch kolejno następujących po sobie latach, zostało z powodzeniem wykorzystane w dalszych programach.

²² Konstrukcja MUP determinuje szybkość wypłat z obligacji. Wykorzystanie danych ze źródeł oficjalnych, umożliwia relatywnie szybką ocenę rozmiaru szkód i wypłat. Najbardziej czasochłonne staje się określenie wypłat z programów, w których wypłaty zależą od portfeli sponsora. W konstrukcji obligacji wyróżnia się tzw. *loss period* – okres, w którym zdarzenie musi nastąpić, by uruchomiono pokrycie – oraz *development period* – okres szacowania szkód. Gdy zdarzenie ma miejsce tuż przed wykupem instrumentu, *moment wykupu* może zostać przesunięty aż do czasu oszacowania szkód i określenia, czy istnieją przesłanki utraty środków przez inwestorów. Szerzej na temat przepływów z obligacji: Małek (2007b, s. 30, tab. 1).

Wykres 1. Struktura utraty kapitału w poszczególnych klasach obligacji wyemitowanych przez Osiris Capital PLC



Źródło: Małek (2007a, s. 18).

Konstrukcję stosowanego indeksu można zatem zdefiniować jako (Blake et al. 2006a, s. 7):

$$indeks_t = \sum_{i=1}^n P_i \sum_{j=1}^{n_j} (\omega_k \cdot G_j \cdot q_{i,j,t}^k + \omega_m \cdot G_j \cdot q_{i,j,t}^m)$$

gdzie:

P_i – waga w indeksie państwa i ,

ω – waga płci w indeksie, przy czym ω_k – waga kobiet, a ω_m – waga mężczyzn

G_j – waga grupy wiekowej j ,

$q_{i,j,t}$ – wskaźnik śmiertelności dla grupy wiekowej j w państwie i w roku t , przy czym $q_{i,j,t}^k$ – wskaźnik dla kobiet, a $q_{i,j,t}^m$ – dla mężczyzn.

Szczególną uwagę należy poświęcić konstrukcji wypłaty dla sponsora. We wszystkich wymienionych powyżej programach utrata środków przez inwestorów, a więc uruchomienie pokrycia, następuje proporcjonalnie po przekroczeniu wartości MUP. Dzięki emisji obligacji sponsor zajmuje długą pozycję w opcji wbudowanej w obligację. W przypadku emisji obligacji w kilku tranzach o różnych stopniach ryzykowności, a więc także różnych momentach uruchamiania zabezpieczenia wypłaty się uzupełniają, tworząc tzw. schemat schodkowy.

Mechanizm strat widać na przykładzie Osiris Capital. Przekroczenie 106% wartości indeksu bazowego (*attachment point*, dolna cena wykonania) uruchamia wypłatę z instrumentu, z najbardziej ryzykownej klasy D, całkowicie pozbawiającej inwestorów wniesionych środków, proporcjonalnie, aż do momentu wzrostu indeksu do 110% (*exhaustion point*, górna cena wykonania). W sytuacji wyższej śmiertelności naturalnym krokiem wynikającym z wyczerpania środków inwestorów klasy D staje się wypłata z klasy C. Po wykorzystaniu pokrycia gwarantowanego w ramach klas C i D, jeśli zachodzi konieczność, cedent otrzyma rekompensatę ze środków klasy B. W przypadku części instrumentów gwarantowanych przez CIFG (por. schemat 2) inwestorzy otrzymają zwrot wartości nominalnej, a roszczenie sponsora spełni gwarant.

Jako uzupełnienie warto zdefiniować proporcjonalną utratę w momencie t :

$$utrata_t = \frac{(\text{indeks}_t - \text{attachment point})}{(\text{exhaustion point} - \text{attachment point})} \cdot 100\%$$

Dolna i górna cena wykonania wyrażona jest zawsze jako wartość indeksu bazowego (np., $1,06 \cdot \text{indeks}_{t=0}$, $1,10 \cdot \text{indeks}_{t=0}$ w przypadku klasy D Osiris). W rzeczywistości może się odnosić więc (w zależności od horyzontu czasowego indeksu) do okresu innego niż roczny. Przykładowo w przypadku programu Axy należy wyodrębnić 3 okresy ekspozycji na ryzyko: 1 stycznia 2006 – 31 grudnia 2007 r.; 1 stycznia 2007 – 31 grudnia 2008 r. i 1 stycznia 2008 – 31 grudnia 2009 r., a wartości cen wykonania szacowane są właśnie w okresach dwuletnich (Standard&Poor's 2006, s. 10).

Jako uzupełnienie rozważań dotyczących wypłat z instrumentów warto przybliżyć rozmiary uruchamiających je zdarzeń. Ogólnie progi uruchamiające pokrycie w analizowanych programach były trudne do osiągnięcia w następstwie pojedynczego zdarzenia, co wynika ze specyfiki zdarzeń katastroficznych.

Według szacunków przytaczanych przez S&P bezwzględne liczby zgonów powodujące utratę kapitału w programie Osiris wynoszą: 898.000 dla klasy B, 641.000 dla inwestorów nabywających obligacje z klasy C i 384.000 – dla klasy D. W przypadku Vita Capital II wzrost indeksu o 10% (uruchamiający pokrycie z klasy D) oznacza, że życie utraciło o 860.000 osób więcej niż w latach 2002–2003. Śmierć dodatkowych 646.000 oraz 430.000 mieszkańców USA, Wielkiej Brytanii, Japonii, Niemiec i Kanady spowoduje płatności z klas, odpowiednio, C i B. Tak wysokie progi byłyby niemożliwe do osiągnięcia nawet w wyniku kumulacji ofiar tsunami z 2004 r., zamachów na WTC i rekordowej epidemii AIDS w 1995 r. (Standard&Poor's 2006, s. 8).

4. Obligacje bazujące na ryzyku długowieczności

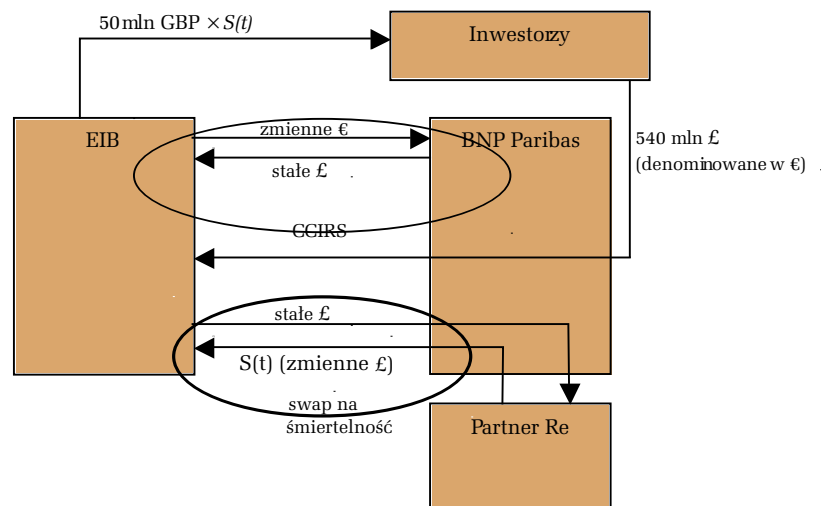
Konstrukcja obligacji opartej na ryzyku długowieczności istotnie różni się od rozwiązań ukształtowanych w wyniku doświadczeń branży majątkowej oraz sekurytyzacji ryzyka wzmożonej śmiertelności, opisanej powyżej. Próba wyemitowania instrumentu zakończyła się niepowodzeniem, co jednak nie przekreśla możliwości ani szans zastosowania tego typu instrumentów. Wyeliminowanie niedociągnięć, które zdecydowały o niepowodzeniu programu, w połączeniu z doświadczeniami z sekurytyzacji pozostałych rodzajów ryzyka może się przyczynić powstania i rozwoju rynku obligacji opartych na ryzyku długowieczności.

4.1. Obligacja EIB/BNP Paribas

Wspomnianą powyżej nieudaną próbę emisji obligacji opartej na ryzyku długowieczności przeprowadził Europejski Bank Inwestycyjny (EIB) we współpracy z BNP Paribas w listopadzie 2004 r. Instrument o wysokości wypłat uwarunkowanych przebiegiem śmiertelności powstał poprzez połączenie klasycznej obligacji o zmiennym oprocentowaniu z dwoma swapami: CCIRS między EIB a BNP Paribas oraz ze swapem opartym na długowieczności, pomiędzy EIB a Partner Re (por. schemat 3) (Blake et al. 2006a, s. 12).

Ryzyko długowieczności miało odzwierciedlać przebieg śmiertelności w ciągu 25 lat (horyzont czasowy instrumentu) mężczyzn z obszaru Walii i Anglii, osiągniętych w 2003 r. 65. rok życia. Uzależnione od niej wartości kolejnych płatności na rzecz inwestorów zostały zdefiniowane jako iloczyn kuponu początkowego (ustalonego jako 50 mln GBP) i indeksu zagregowanej śmiertelności populacji $S(t)$ danej wzorem (Blake et al. 2006a, s. 10; 2006b, s. 5–6):

Schemat 3. Struktura programu EIB/BNP Paribas



Źródło: opracowanie własne na podstawie: Blake et al. (2006a, s. 11).

$$S(t) = S(0) \times (1 - m(2003; 65)) \times (1 - (m(2004; 66) \times (1 - m((2002 + t); (64 + t)))) \quad (3)$$

gdzie $m(y, x)$ jest współczynnikiem śmiertelności osób w wieku x opublikowanym w roku y przez The Office for National Statistics.

Przejrzystość informacyjna wynikająca z wykorzystania oficjalnych danych bezspornie zwiększała atrakcyjność instrumentu. Sam indeks, niezdywersyfikowany pod względem wieku ani płci i dodatkowo opierający się na „czystych” stopach śmiertelności, okazał się jednak niewystarczająco reprezentatywny, by być narzędziem efektywnego hedgingu dla większej liczby inwestorów. Dla podmiotów mających portfel o nieco odmiennej strukturze wiekowej czy ze znacznym udziałem kobiet ewentualne ryzyko bazy wynikające z takiego zabezpieczenia było zbyt wysokie. Ryzyko bazy wiązano również ze sztywno ustanowionym poziomem płatności emerytalnych. W programie nie został w ogóle uwzględniony fakt, że w rzeczywistości wypłacane świadczenia są waloryzowane, co ma szczególne znaczenie ze względu na dwudziestopięcioletni horyzont wypłat (Blake et al. 2006a, s. 10–14). Jeśli chodzi o ryzyko bazy, należy podnieść zarzut także uwzględnienia „łysych” stóp śmiertelności. Taka konstrukcja mogła spowodować, że przebieg śmiertelności w portfelach konkretnych inwestorów będzie inny niż w populacji bazowej (dotyczy to zwłaszcza podmiotów o mniejszym udziale w rynku) (Azzopardi 2005, slajd nr 8).

Jak wspomniano wcześniej, ustalenie MUP jest niezwykle trudne. Przykład EIB dowodzi jednak, że opieranie się na „łysych” stopach śmiertelności generuje

zbyt duże ryzyko bazy, by emisja mogła zakończyć się sukcesem. O ile w przypadku długowieczności niemożliwie jest skonstruowanie MUP parametrycznych, o tyle pod rozważę powinny być wzięte warianty MUP uzależnionych od wartości indeksów branżowych czy szkód wynikających z wartości modelu. Możliwe jest też rozwiązanie łączące tablice życia dla kobiet i mężczyzn, w różnym wieku, z odpowiednio dobranymi wagami – czyli swoista klasyczna dywersyfikacja indeksu.

4.2. Obligacje na ryzyko długowieczności

Przed przedstawieniem konkretnych rozwiązań dotyczących obligacji należy zauważyć, że może ona przyjąć formę zabezpieczenia „na długowieczność” (*longevity bond*) lub „na przeżycie” (*survivor bond*) konkretnej populacji. Pierwsze podejście uzależnia wypłatę dla inwestorów od przebiegu śmiertelności w populacji w konkretnym okresie, z reguły zbieżnym z horyzontem emisji instrumentu. Drugi typ generuje przepływy pieniężne do momentu pozostawania przy życiu ostatniego członka danej populacji (Blake et al. 2006a, s. 16).

Wariant drugi rozwiązuje problem zabezpieczenia niespójnego z horyzontem czasowym produktów wchodzących w skład portfela inwestorów – w przypadku zastosowanego przez Europejski Bank Inwestycyjny schematu *longevity* zapadalność instrumentu, zbieżna z ukończeniem przez członków populacji 89. roku życia, przerywała zabezpieczenie w stosunku do wszystkich osób przekraczających ten wiek, pozostawiając dalsze ryzyko wypłat podmiotowi nabywającemu obligację. Innym sposobem rozwiązania problemu jest przesunięcie

terminów płatności (tzw. *deferred longevity bond*). Ma ono dodatkową zaletę: że rozwiązuje problem wypłat w pierwszych latach, gdy ryzyko jest uznawane za niewielkie, a zabezpieczenie za zbędne²³.

Dzięki zastosowaniu wariantu *survivor* można wyeliminować jeden z przejawów ryzyka bazy – wynikający z faktu, że portfele inwestorów mogą zawierać produkty ubezpieczeniowe wypłacające świadczenia osobom przeżywającym osoby wchodzące w skład populacji indeksowej czy, w przypadku rent okresowych, renty o dłuższym horyzoncie wypłaty niż okres wykupu. Problem dopasowania zabezpieczenia do specyfiki portfeli inwestorów mógłby być rozwiązany poprzez emisję obligacji zerokuponowych o różnych populacjach indeksowych i (lub) terminach wykupów (Blake et al. 2006a, s. 13, 16), a nawet – zdaniem autorki – opierających się na różnych typach indeksów i MUP. Na odpowiednio płynnym rynku kombinacja instrumentów tego typu powinna umożliwić skonstruowanie portfela zabezpieczeń dopasowanego do skali i specyfiki prowadzonej działalności.

Jeśli chodzi o strukturyzowanie instrumentu, to zdaniem autorki warto przemyśleć zastosowanie konstrukcji z zagrożoną wartością nominalną (*principal at risk bond*), atrakcyjnej z co najmniej trzech powodów. Po pierwsze, emisja obligacji tego typu pozwala na uzyskanie przez sponsora większego pokrycia, dzięki czemu może być zmniejszony wolumen emisji, co na tworzącym się dopiero rynku sekurytyzacji długowieczności, o ograniczonej liczbie czy niewielu grupach potencjalnych inwestorów mogłoby otworzyć drogę nowym programom. Takie podejście umożliwia różnicowanie warunków utraty kapitału poprzez oferowanie odmiennych pod względem ryzykowności klas czy transz, i w konsekwencji dotarcie do inwestorów o różnych preferencjach co do ryzyka. Rozwiązanie to sprawdziło się w przypadku wielu rodzajów ryzyka: nie tylko katastrofowych, ale także terrorystycznych czy wzmożonej śmiertelności.

Interesujące wydaje się zastosowanie schematu sekurytyzacji wtórnej, przedstawionej w rozdziale drugim, do ryzyka długowieczności. W tym przypadku ryzyko transferowane jest od sponsora, wraz ze składką π (formalnie z tytułu ochrony reasekuracyjnej), do spółki specjalnego przeznaczenia (SPV), która emituje obligacje P i inwestuje otrzymane z tego tytułu środki w bezpieczne papiery wartościowe. Ujemnymi przepływami pieniężnymi podmiotu są kupony (D_t) i wypłaty dla sponsora (B_t), których wysokość zależy od wartości indeksu.

Uwzględniając wariant wypłat na podstawie opcji *call spread*, można – przy założeniu stałej²⁴ wartości renty: 1000 jednostek pieniężnych dla wszystkich członków populacji indeksowej – przedstawić strukturę B_t jako (Lin, Cox 2005, s. 229):

$$B_t = \begin{cases} 1000C & \text{dla } I_{x+t} > C + X_t \\ 1000(I_{x+t} - X_t) & \text{dla } X_t < I_{x+t} \leq C + X_t \\ 0 & \text{dla } I_{x+t} \leq X_t \end{cases}$$

gdzie:

I_{x+t} – liczba osób z populacji dożywających roku t (osoby w wieku $(x + t)$),

x – wiek „początkowy” populacji (w roku $t = 0$),

X_t – MUP określony jako liczba osób dożywających wieku $(x + t)$,

C – maksymalne pokrycie z tytułu obligacji (jako liczba osób).

Od wartości tego samego indeksu zależy wysokość kuponów D_t (Lin, Cox 2005, s. 229):

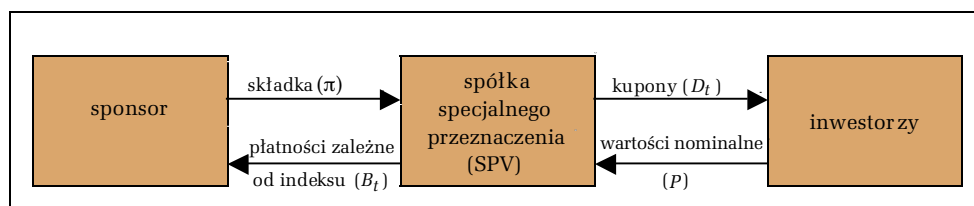
$$D_t = \begin{cases} 0 & \text{dla } I_{x+t} > C + X_t \\ 1000(C + X_t - I_{x+t}) & \text{dla } X_t < I_{x+t} \leq C + X_t \\ 1000C & \text{dla } I_{x+t} \leq X_t \end{cases}$$

Nietrudno zatem zauważyć, że suma kuponów i płatności dla sponsora, czyli przepływów ujemnych

²³ Por. Blake et al. (2006a).

²⁴ W rzeczywistości renta powinna być jednak waloryzowana.

Schemat 4. Struktura sekurytyzacji ryzyka długowieczności w wariacie pośrednim



SPV, w każdym momencie t równa jest wartości maksymalnego pokrycia: $D_t + B_t = 1000C$.

5. Podsumowanie

Instrumenty bazujące na ryzyku śmiertelności budzą wiele kontrowersji. Trwają spory o konstrukcję instrumentów, podmiotów emitujących czy możliwości wykorzystania sekurytyzacji w programach zabezpieczenia społecznego (MacMinn et al. 2006).

Doświadczenia branży majątkowej i udane programy transferujące na rynek kapitałowy ryzyko ekstremalnej śmiertelności dowodzą, że sekurytyzacja jest atrakcyjnym rozwiązaniem dla branży asekuracyjnej i należy dostrzegać jej potencjał (Swiss Reinsurance Group 2006). Zdaniem autorki za rozwojem rynku przemawia także możliwość wykorzystania naturalnego hedgingu na płaszczyźnie ryzyka: ekstremalnej śmiertelności i długowieczności – do tego stopnia, że wydaje się, że potencjalny sponsor obligacji na wzmogłą śmiertelność może inwestować w obligacje na długowieczność i odwrotnie.

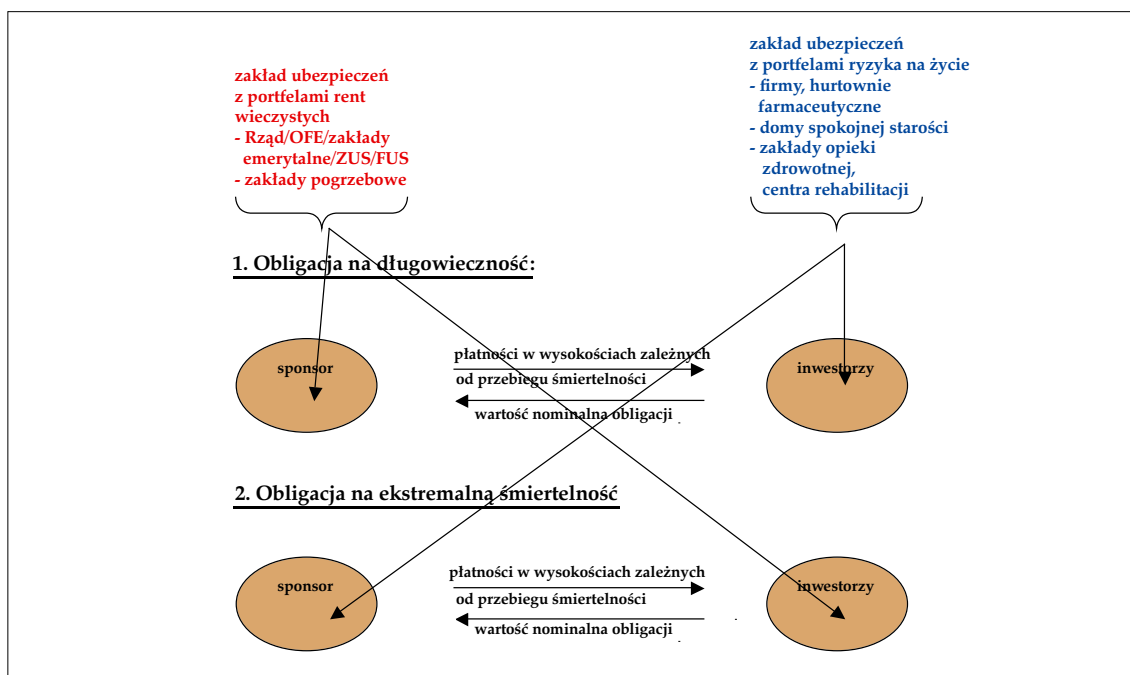
Przy odpowiednim ustrukturyzowaniu instrumentów ubezpieczyciele z portfelami terminowymi na życie oraz podmioty wypłacające renty i emerytury (głównie wieczyste) mogą stworzyć podstawy nowego – efektywnego i atrakcyjnego – sposobu cedowania ryzyka śmiertelności. Oprócz uczestników rynku asekuracyjnego

jest bowiem wiele podmiotów spoza branży również narażonych na to ryzyko: firmy i hurtownie farmaceutyczne, instytucje świadczące usługi zdrowotne i rehabilitacyjne, domy spokojnej starości czy zakłady pogrzebowe. Umiejętne powiązanie interesów podmiotów zabezpieczających się przed zbyt wysoką umieralnością i tych, dla których zagrożeniem jest wydłużające się trwanie życia, powinno zwiększyć zainteresowanie alternatywnym zabezpieczeniem.

Ciekawe wydaje się zastosowanie obligacji opartej na długowieczności w społecznych programach zabezpieczenia na starość. Podmiot zobowiązany do zaspokajania roszczeń mógłby być zarówno sponsorem, jak inwestorem (zakładając istnienie odpowiednich regulacji prawnych). Mechanizm zabezpieczenia przedstawia schemat 5. Jeśli celem podmiotu sponsorującego jest zabezpieczenie przed ryzykiem wzrostu długości życia i (lub) spadku śmiertelności, to sensowne okazuje się zainwestowanie w taki instrument przez podmiot, który ponosi straty w wyniku skrócenia trwania życia i (lub) wzrostu śmiertelności. Wysoka śmiertelność gwarantuje w tym przypadku²⁵ wyższe płatności z tytułu obligacji, więc mogą one stanowić rekompensatę poniesionych strat. Analogicznie – w przypadku obligacji na wzmogłą śmiertelność wyższy kupon, wynikający z niskiej śmiertelności, rekompensuje straty i wyższe zobowiązania inwestorów.

²⁵ Zakładając, że konstrukcja indeksu jest zgodna ze wzorem (3).

Schemat 5. Podmioty sekurytyzacji ryzyka śmiertelności



Konstrukcja zilustrowana na schemacie 5 wydaje się prosta, należy jednak stwierdzić, że w rzeczywistości sama emisja obligacji generuje znaczne koszty i jest opłacalna jedynie przy wysokim wolumenie (Blake et al. 2006a, s. 19). Ponadto, uzależnienie płatności od przebiegu śmiertelności znacznie komplikuje przepływy (pomocne okazuje się wykorzystanie mechanizmu swapowego). Z tego względu może się okazać, że na kształtującym się rynku pierwszeństwo będą miały swapy.

Wraz z rozwojem rynku obligacje transferujące ryzyko śmiertelności powinny jednak zyskiwać na znaczeniu. Warunkiem jest ich odpowiednie ustrukturyzowanie, zapewniające atrakcyjność i zainteresowanie ze strony potencjalnych graczy rynku większe niż strach przed nowością²⁶.

²⁶ G. Tett i J. Chung (2007) stawiają tezę, że winę za niepowodzenie obligacji EIB ponosi ich innowacyjność i wynikająca stąd nieświadomość inwestorów, jakie korzyści może przynieść ten instrument.

Bibliografia

- Albrecht P., Schradin H. (1998), *Alternativer Risikotransfer: Verbriefung von Versicherungsrisiken*, „Zeitschrift für die gesamte Versicherungswissenschaft“, Band 87, s. 573–610.
- Axa (2006), *Axa announces the successful completion of its first mortality risk securitization transaction*, Press Release, 13 November, http://www.axa.com/lib/en/uploads/pr/group/2006/AXA_PR_20061113.pdf
- Azzopardi M. (2005), *The longevity bond*, prezentacja na konferencji “First International Conference on Longevity Risk and Capital Market Solutions”, Pensions Institute, American Risk and Insurance Association, Centre for Risk and Insurance, Nottingham University Business School, 18 February, London, http://www.pensions-institute.org/conferences/longevity/Azzopardi_Mark.pdf
- Bauer D., Kramer F. (2008), *Risk and Valuation of Mortality Contingent Catastrophe Bonds*, “Discussion Paper”, No. 08–05, Pensions Institute, London.
- Blake D., Cairns A., Dowd K. (2006a), *Living with mortality: Longevity bonds and other mortality-linked securities*, mimeo, Institute of Actuaries and Faculty of Actuaries, <http://www.ma.hw.ac.uk/~andrewc/papers/baj2006.pdf>
- Blake D., Cairns A., Dowd K., MacMinn R. (2006b), *Longevity Bonds: Financial Engineering, Valuation and Hedging*, “Working Paper”, No. 0617, Pensions Institute, London.
- Bütow S. (2001), *Securitisation in der Personenrückversicherung*, „Perspektiven der Hannover Rück zu aktuellen Themen der internationalen Lebensversicherung“, Ausgabe Nr 7, Hannover.
- Cowley A., Cummins D. (2005), *Securitization in Life Insurance Assets and Liabilities*, “Journal of Risk and Insurance”, Vol. 72, No. 2, s. 193–226.
- Dubinski W., Laster D. (2003), *Insurance-linked Securities*, mimeo, Swiss Re Capital Markets Corporation, New York.
- Hartung T. (2004), *Alternative Risikotransfer-Instrumente*, „Wirtschaftswissenschaftliches Studium (WiSt)“, Juli, s. 401–405.
- Lin Y., Cox S. (2005), *Securitization of mortality risks in life annuities*, “Journal of Risk and Insurance”, Vol. 71, No. 2, s. 227–252.
- MacMinn R., Brockett P., Blake D. (2006), *Longevity Risk and Capital Markets*, “Journal of Risk and Insurance”, Vol. 73, No. 4, s. 551–557.
- Małek A. (2007a), *Jak sekurytyzować życie?*, „Gazeta Bankowa”, nr 3, s. 15–19.
- Małek A. (2007b), *Sekurytyzacja ryzyka terroryzmu*, „Wiadomości Ubezpieczeniowe”, nr 5–6/2007, s. 28–36.
- Monday B. (2006), *“Don’t count your chickens because they’ll scratch”- The threat of a global pandemic including Bird Flu*, “Hannover Re’s Perspectives, Current Topics of International Life Insurance”, http://www.hannover-re.com/resources/hlr/generic/hlr/publications/Schriftenreihe_Nr12_engl.pdf
- PricewaterhouseCoopers (2006), *Innovative financing: Life insurance securitization*, [http://www.pwc.com/extweb/pwcpublications.nsf/docid/D26E84E618B31A5B852570FA0055E4FC/\\$File/lisecuritisation.pdf](http://www.pwc.com/extweb/pwcpublications.nsf/docid/D26E84E618B31A5B852570FA0055E4FC/$File/lisecuritisation.pdf)
- Punter A. (2000), *The changing Economics of Non-life Insurance: New Solutions for the Financing of Risk*, referat na konferencję “UK Insurance Economists’ Conference”, Centre for Risk & Insurance Studies, The University of Nottingham, 29–30 March, Nottingham.
- Ronka-Chmielowiec W. (red.) (2004), *Zastosowanie metod ekonometryczno-statystycznych w zarządzaniu finansami w zakładach ubezpieczeń*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- Reactions (2006), *Scottish Re mortality bond hits target despite pandemic fears*, 9 May, <http://www.reactionsnet.com/default.asp?Page=3&ISS=21876>

- Scottish Re (2006), *Scottish Re Closes \$155 Million Mortality Catastrophe Bond*, Press Release, 5 April, Hamilton.
- Standard&Poor's (2003), *Presale: Vita Capital Ltd.'s Principal-At-Risk Variable-Rate Mortality Catastrophe-Indexed Note*, New York.
- Standard&Poor's (2006), *Osiris Capital PLC. Commentary Report*, New York.
- Stone Ch., Zissu A. (2006), *Securitization of Senior Life Settlements: Capturing Value from Early Death*, <http://www.mortalityrisk.org/Papers/Settlements/LifeSettlementsJOD2006.pdf>
- Strube M. (2001), *Alternativer Risikotransfer von Katastrophenrisiken: die Rückversicherung mit Anleihen und börsengehandelten Optionen im Vergleich*, Deutscher Universitäts-Verlag, Wiesbaden.
- Swiss Re (1996), *Insurance derivatives and securitization: New hedging perspectives for US catastrophe insurance market?*, "Sigma", No 5, Zurich.
- Swiss Re (2003), *Swiss Re obtains USD 400 million of extreme mortality coverage- its first life securitization*, News Release, 8 August, Zurich.
- Swiss Re (2005), *Swiss Re successfully closes its second life catastrophe bond and obtains USD 362 million of mortality risk coverage through the Vita Capital II programme*, News Release, 14 April, Zurich.
- Swiss Re (2006), *Securitization – new opportunities for insurers and investors*, "Sigma", No 7/2006, Zurich.
- Swiss Re (2007), *Swiss Re obtains USD 705 Million of Extreme Mortality Risk Protection through its Vita Capital programme*, News Release, 16 January, Zurich.
- Tett G., Chung J. (2007), *Death and the salesman*, "Financial Times UK", 24.02.2007, <http://www.ft.com/cms/s/f7a24c26-c3ac-11db-9047-000b5df10621.html>
- Wagner F. (1997), *Risk Securitization als alternatives Mittel des Risikotransfers von Versicherungsunternehmen*, "Zeitschrift für die gesamte Versicherungswissenschaft", No. 86, s. 510–552.

Partnerstwo publiczno-prywatne – dlaczego warto zabiegać o jego rozwój

Public-Private Partnership – Why Should We Insist on its Development?

Agnieszka Cenkier*

pierwsza wersja: 13 lutego 2008 r., ostateczna wersja: 16 kwietnia 2008 r., akceptacja: 23 kwietnia 2008 r.

Streszczenie

Partnerstwo publiczno-prywatne, jako jedna z metod wykonywania zadań publicznych, wiąże się z możliwością uzyskania w wyniku realizacji danego przedsięwzięcia korzyści dla interesu publicznego przewyższających korzyści możliwe do uzyskania w przypadku realizacji tego przedsięwzięcia w inny sposób. Źródłem tych korzyści jest udział w partnerstwie publiczno-prywatnym sektora prywatnego – z jego kapitałem, wiedzą, umiejętnościami oraz doświadczeniem. Warunkiem koniecznym ich uzyskania jest odpowiedni (optymalny) podział ryzyka przedsięwzięcia między sektor publiczny a inwestorów prywatnych, zgodny z ich przygotowaniem do kontrolowania określonych rodzajów ryzyka. *Value for money* pozwala zmierzyć korzyści wynikające z realizacji projektu w trybie partnerstwa publiczno-prywatnego i ocenić zasadność jego zastosowania.

Słowa kluczowe: partnerstwo publiczno-prywatne, podział ryzyka w projektach partnerstwa publiczno-prywatnego, *value for money*

Abstract

Public-private partnership is a delivery model for new infrastructure to meet social needs that should only be used where it can be demonstrated that it offers greater value for money for the public sector compared with other forms of procurement. Private sector should only be involved in building and maintenance of public infrastructure if it has such expertise to deliver that can be expected to offer value for money. The appropriate sharing of risk is the key to ensuring that value for money benefits in public-private partnerships are indeed obtained. Transferring risks to the private sector is not an end in itself but it has to force private investors to create a better outcome. It is a key issue to ensure value for money performance in public-private partnership.

Keywords: public-private partnership, risk sharing in PPP projects, value for money

JEL: R53

* Szkoła Główna Handlowa w Warszawie, Kolegium Zarządzania i Finansów, Katedra Finansów Przedsiębiorstwa, e-mail: acenkier@sgh.waw.pl

1. Wstęp

Podstawowym celem artykułu jest wstępne zapoznanie Czytelnika z istotą partnerstwa publiczno-prywatnego jako metody wykonywania zadań publicznych oraz zaprezentowanie korzyści, które można uzyskać dzięki jego stosowaniu. Opierając się na studiach literatury przedmiotu oraz praktycznych doświadczeniach, przede wszystkim zagranicznych (np. British Columbia Ministry of Municipal Affairs 1999; Komisja Europejska 2003; European Commission 2004; Moszoro 2005; HM Treasury 2006), w artykule podjęto próbę zaprezentowania partnerstwa publiczno-prywatnego. W szczególności skoncentrowano się na tych właściwościach, które – jak się wydaje – stanowią o jego istocie.

Prezentowane w opracowaniu cechy partnerstwa publiczno-prywatnego zostały sformułowane głównie na podstawie doświadczeń Wielkiej Brytanii i Kanady. W państwach tych partnerstwo publiczno-prywatne jest stosowane na największą skalę pod względem liczby i wartości przedsięwzięć zrealizowanych w tym trybie (np. British Columbia Ministry of Municipal Affairs 1999; Allen 2003; Eaton, Akbiyikli 2005). W artykule wykorzystano materiały na temat funkcjonowania partnerstwa publiczno-prywatnego, publikowane przez administrację szczebla centralnego tych państw. Wśród tych opracowań można wyróżnić publikacje stanowiące rodzaj podręczników (przewodników) partnerstwa publiczno-prywatnego. Mają one ułatwić podmiotom publicznym zainteresowanym partnerstwem publiczno-prywatnym zrozumienie jego istoty oraz praktyczne wykorzystanie do realizacji zadań publicznych (np. British Columbia Ministry of Municipal Affairs 1999; The Stationery Office 2000; Komisja Europejska 2003; PPP Unit 2003; HM Treasury 2007). Są również opracowania podsumowujące dotychczasowe doświadczenia w stosowaniu partnerstwa publiczno-prywatnego, których analiza ma służyć usprawnieniu (np. Bates 1997; NAO 1998a; 1998b; State Government of Victoria Department of Treasury and Finance 2001; European Commission 2004; Eaton, Akbiyikli 2005; Deloitte 2006b; CBI 2007a; CBI 2007b). Wyniki badań przeprowadzonych dotychczas przez autorkę na podstawie dostępnych źródeł wydają się wskazywać, że doświadczenia różnych państw, którym udało się upowszechnić partnerstwo publiczno-prywatne w mniejszym lub większym stopniu, są bardzo podobne (Deloitte 2006a). Bardzo ograniczona jest natomiast możliwość wykorzystania doświadczeń polskich, ze względu na ich małą skalę.

Partnerstwo publiczno-prywatne charakteryzuje się przede wszystkim wspólnym, długookresowym zaangażowaniem sektora publicznego i sektora prywatnego w realizację zadań publicznych. Ryzyko danego przedsięwzięcia zostaje rozłożone pomiędzy oba sektory w sposób pozwalający każdej ze stron przyjąć to ryzyko, do zarządzania którym jest lepiej przygotowana (HM

Treasury 2003, s. 35-37; Moszoro 2005, s. 47-49). O partnerstwie publiczno-prywatnym można mówić w szerszym lub węższym znaczeniu. W szerszym ujęciu partnerstwem publiczno-prywatnym jest w zasadzie każda długookresowa forma współpracy sektora publicznego z sektorem prywatnym, która służy zaspokojeniu potrzeb społecznych. Podejście to reprezentuje na przykład Komisja Europejska (2003). W węższym znaczeniu do partnerstwa publiczno-prywatnego zalicza się jedynie te przypadki współpracy sektora publicznego i prywatnego w zakresie inwestycji publicznych, które spełniają określone warunki. Takie węższe ujęcie partnerstwa publiczno-prywatnego zaproponował polski ustawodawca. Ustawa o partnerstwie publiczno-prywatnym ogranicza je do tych przedsięwzięć, które są realizowane na zasadach określonych w ustawie¹.

Autorce bliższe jest podejście wynikające z szerszego rozumienia partnerstwa publiczno-prywatnego. Zgodnie z nim **teoria partnerstwa publiczno-prywatnego powstaje na podstawie empirycznych doświadczeń i nie ogranicza jego ewolucji i rozwoju, których wyrazem jest np. różnorodność stosowanych rozwiązań.**

Wykorzystanie partnerstwa publiczno-prywatnego do wykonywania zadań publicznych prowadzi do wzrostu liczby realizowanych projektów, co w wyniku z doświadczeń państw, w których partnerstwo jest stosowane. Ponadto pozwala uzyskać korzyści dla interesu publicznego, których osiągnięcie nie jest możliwe w przypadku realizacji tego typu zadań w sposób tradycyjny (UNIDO 2006). Zgodnie z ustawą, korzyścią dla interesu publicznego mogą być w szczególności: zmniejszenie wydatków sektora publicznego, podwyższanie standardu oferowanych usług lub obniżenie uciążliwości dla otoczenia. Ostatecznym beneficjentem osiąganych korzyści jest społeczeństwo, które zyskuje dostęp do usług publicznych o wysokiej jakości, świadczonych za pomocą nowoczesnej infrastruktury (Deloitte 2006a).

Współpraca podejmowana w ramach partnerstwa publiczno-prywatnego przez podmioty publiczne i przedsiębiorców prywatnych, pozwala każdej ze stron osiągnąć zamierzony cel. W przypadku sektora publicznego celem tym jest zaspokajanie potrzeb zbiorowych, podczas gdy inwestorzy prywatni podejmują współpracę kierowani chęcią uzyskania satysfakcjonującej ich stopy zwrotu. Z punktu widzenia inwestorów prywatnych projekty partnerstwa publiczno-prywatnego są (jednymi z wielu) okazjami inwestycyjnymi, związanymi z możliwością uzyskania określonej stopy zwrotu, charakteryzującymi się określonym poziomem ryzyka. Podmioty prywatne uczestniczą w ich realizacji z pobudek ekonomicznych. Inwestorzy prywatni są zainteresowani udziałem tylko w takich przedsięwzięciach, które zapewnią im osiągnięcie satysfakcjonującego zys-

¹ Ustawa z dnia 28 lipca 2005 r. o partnerstwie publiczno-prywatnym, Dz.U. nr 169, poz. 1420, zob. też: *Partnerstwo publiczno-prywatne z praktycznym komentarzem, stan prawny na 19 listopada 2007 r.*, „Gazeta Prawna”, INFOR, Warszawa.

ku. Angażują w wykonanie danego przedsięwzięcia dostępne kapitały, ale również swoją wiedzę, umiejętności i doświadczenie.

Jednoczesna realizacja, w ramach jednego przedsięwzięcia, celów społecznych (sektora publicznego) oraz komercyjnych (sektora prywatnego) jest możliwa dzięki podziałowi między oba sektory zadań związanych z realizacją projektu w sposób umożliwiający optymalną alokację właściwego mu ryzyka. Różnice między rozwiązaniami dotyczącymi podziału zadań i ryzyka w odniesieniu do poszczególnych przedsięwzięć, są przyczyną występowania różnych wariantów partnerstwa publiczno-prywatnego, często nazywanych modelami partnerstwa publiczno-prywatnego. Modele partnerstwa publiczno-prywatnego (na przykład BOT – *Build-Operate-Transfer* czy DFBO – *Design-Finance-Build-Operate*), prezentowane w publikacjach podejmujących tę tematykę, krajowych (niezbyt licznych) i zagranicznych, stanowią katalog otwarty. Powstał on głównie na podstawie dotychczasowych doświadczeń i nie wyklucza modyfikacji istniejących rozwiązań stosownie do specyfiki realizowanych przedsięwzięć (Komisja Europejska 2003; Korbus, Strawiński 2006)².

Rozwój partnerstwa publiczno-prywatnego obserwowany w ostatnich kilkunastu latach w wielu państwach początkowo był szczególnie widoczny w państwach zachodnioeuropejskich. Obecnie wydaje się przybierać na sile i zataczać coraz szersze kręgi. Choć zaawansowanie tego procesu w poszczególnych państwach jest zróżnicowane, jego przesłanki są zbliżone (Deloitte 2006a).

Struktura opracowania jest następująca. Przesłanki rozwoju partnerstwa publiczno-prywatnego zostały zaprezentowane w rozdziale 2.

Na treść rozdziału 3. składa się charakterystyka korzyści, możliwych do uzyskania w wyniku stosowania partnerstwa publiczno-prywatnego, ze szczególnym uwzględnieniem źródeł tych korzyści.

Przedmiotem rozważań w rozdziale 4. jest ryzyko w projektach partnerstwa publiczno-prywatnego. W rozdziale tym zostały przedstawione różne rodzaje ryzyka, ze wskazaniem przykładowych klasyfikacji. Omówiono też kwestię podziału ryzyka projektu między strony zaangażowane w jego wykonanie, decydującą dla pomyślnej realizacji przedsięwzięcia, której wyrazem są odnoszone korzyści.

Rozdział 5. zawiera krótką charakterystykę instrumentów stosowanych do oceny korzyści z realizacji przedsięwzięcia na zasadach partnerstwa publiczno-prywatnego.

2. Przesłanki rozwoju partnerstwa publiczno-prywatnego

Podmioty publiczne są zobowiązane do finansowania zadań publicznych. W tym celu wyposażono je w środki publiczne. Jednak **potrzeby w zakresie rozwoju i modernizacji infrastruktury oraz świadczonych za jej pomocą usług publicznych w poszczególnych państwach z reguły znacznie przewyższają możliwości ich finansowania przez sektor publiczny**. Brak równowagi między skalą potrzeb i możliwościami ich finansowania jest przyczyną występowania właściwie we wszystkich krajach chronicznych zaległości w realizowaniu przez administrację publiczną zadań publicznych, służących zaspokajaniu potrzeb zbiorowych. Ten **stan niedoinwestowania utrzymujący się w dłuższym okresie prowadzi do sytuacji, w której – ze względu na zaniedbaną, przestarzałą infrastrukturę – nie jest możliwe świadczenie usług publicznych o standardzie oczekiwanym przez społeczność poszczególnych regionów**. Dzięki rozwojowi nowych technologii i postępującej globalizacji świadczenia publiczne o wysokim standardzie są dostępne na coraz większą skalę. Tam, gdzie jeszcze nie dotarły, są niecierpliwie oczekiwane przez mieszkańców (Amerykańska Izba Handlowa w Polsce 2002; Deloitte 2006a). Rzecz jasna, skala zaniedbań i wywołanych nimi konsekwencji oraz poziom oczekiwań społecznych nie są jednakowe we wszystkich państwach. Zróżnicowanie to jest wynikiem m.in. tradycji zaspokajania społecznych potrzeb i osiągniętego poziomu rozwoju gospodarczego poszczególnych krajów.

W krajach zachodnioeuropejskich zaniedbania w rozwoju infrastruktury zostały spotęgowane wyraźną od lat 70. ubiegłego stulecia tendencją do ograniczania wydatków publicznych na realizację zadań publicznych. Prowadzona przez rządy tych państw polityka w zakresie inwestycji publicznych, wynikająca w znacznym stopniu ze złej kondycji ich finansów publicznych, relatywnie szybko doprowadziła infrastrukturę do stanu wymagającego natychmiastowej odbudowy, a następnie rozwoju (Eaton, Abaiyikli 2005).

Czynnikiem dodatkowo dyscyplinującym napięte budżety finansów publicznych wielu państw europejskich, a tym samym dodatkowo utrudniającym finansowanie zadań publicznych ze środków publicznych, okazało się opublikowanie w Traktacie z Maastricht kryteriów zbieżności, których spełnienie stało się zadaniem priorytetowym. Starania rządów, zmierzające do utrzymania deficytu i długu sektora finansów publicznych na poziomie nieprzekraczającym wartości referencyjnych, istotnie ograniczyły zdolność sektora publicznego do podejmowania inwestycji, mających na celu doprowadzenie infrastruktury do stanu umożliwiającego świadczenie usług publicznych o odpowiednio wysokim standardzie, finansowanych wyłącznie ze środków publicznych. W tej sytuacji wykonywanie zadań publicznych

² Ze względu na wstępny i ogólny charakter opracowania w odniesieniu do wielu zagadnień autorka ogranicza się jedynie do ich zasygnalizowania. Okazją do ich rozwinięcia mogą być kolejne artykuły powstające w ramach cyklu poświęconego problematyce partnerstwa publiczno-prywatnego. BOT i DFBO są jednymi z podstawowych, powszechnie stosowanych modeli partnerstwa publiczno-prywatnego.

stało się szczególnie trudne, zwłaszcza że w tym samym czasie nasiliły się oczekiwania społeczne w zakresie dostępności i jakości usług publicznych świadczonych przez administrację publiczną (np. HM Treasury 2003; PPP Unit 2003).

Unowocześnienie i rozbudowa infrastruktury w stopniu umożliwiającym zaspokajanie potrzeb społecznych na odpowiednio wysokim, oczekiwanym poziomie, z wykorzystaniem dotychczasowych metod wykonywania zadań publicznych były trudne właściwie dla wszystkich zainteresowanych państw, zwłaszcza w warunkach zaostrenia dyscypliny finansów publicznych. Wobec powyższego naturalne wydaje się **zainteresowanie rządów możliwościami wykonywania zadań publicznych, służących zaspokajaniu potrzeb społecznych, przy współudziale inwestorów prywatnych**. Nic dziwnego, że właśnie w partnerstwie publiczno-prywatnym dostrzeżono szansę na relatywnie szybkie rozwiązanie problemu. Ostatnia dekada ubiegłego wieku to czas, w którym do realizacji inwestycji publicznych zaczęto stosować partnerstwo publiczno-prywatne. Moment wystąpienia zjawiska (początek lat 90.) oraz jego natężenie wydają się wskazywać, że zainteresowanie sektora publicznego współpracą z sektorem prywatnym wiąże się z koniecznością spełnienia fiskalnych warunków zbieżności z Maastricht.

W początkowym okresie rozwoju partnerstwa publiczno-prywatnego, w sytuacji dotkliwego niedoboru środków publicznych, na którego rychłe przewyciężenie raczej się nie zanosilo, dla podmiotów publicznych szczególnie interesująca wydawała się możliwość finansowania zadań publicznych kapitałem inwestorów prywatnych, niedostępnym bez ich zaangażowania. **Możliwość realizacji przez sektor publiczny zadań publicznych bez konieczności wydatkowania środków publicznych, w trybie niezwiązanym z cyklem budżetu finansów publicznych, musiała przyciągnąć uwagę rządów wielu państw**. Już pierwsze próby realizacji inwestycji na zasadach partnerstwa publiczno-prywatnego pokazały, że taka współpraca nie tylko pozwala zaoszczędzić publiczne pieniądze, ale może również przynieść inne korzyści wynikające z zaangażowania inwestorów prywatnych w wykonanie zadań publicznych³. Właśnie w możliwych korzyściach ze stosowania partnerstwa publiczno-prywatnego należy widzieć podstawową przyczynę jego rozwoju, zapoczątkowanego w latach 90. i, z różnym nasileniem, ciągle obserwowanego w wielu krajach (HM Treasury 2003).

3. Zagadnienie korzyści w partnerstwie publiczno-prywatnym

Partnerstwo publiczno-prywatne w żadnym wypadku nie jest sposobem na rozwiązanie wszelkich problemów

związanych z realizacją zadań publicznych przez podmioty publiczne. **Jest to jedna z metod, które mogą być zastosowane w przedsięwzięciach podejmowanych przez sektor publiczny w zakresie zaspokajania potrzeb społecznych**. Przewaga partnerstwa publiczno-prywatnego jako metody wykonywania zadań publicznych polega na tym, że pozwala ono uzyskać dla interesu publicznego większe korzyści niż w przypadku realizacji tego przedsięwzięcia w inny sposób (HM Treasury 2004).

3.1. Korzyści dla interesu publicznego jako główne kryterium wyboru sposobu wykonywania zadań publicznych

Podstawową zasadą, którą kierują się podmioty publiczne przy wyborze metody wykonywania poszczególnych zadań publicznych, jest uzyskanie przez sektor publiczny jak największych korzyści (*value for money*⁴). **Partnerstwo publiczno-prywatne powinno być wybierane tylko wtedy, kiedy prowadzi do uzyskania większych korzyści niż w wyniku zastosowania innych metod realizacji danego przedsięwzięcia. W przeciwnym przypadku projekt powinien zostać zrealizowany metodą tradycyjną**.

Kryterium korzyści jest powszechnie stosowane w wielu krajach i często jest obecne w przepisach regulujących funkcjonowanie partnerstwa publiczno-prywatnego. Na przykład Komisja Europejska (2003, wstęp) wyraża pogląd, że partnerstwo publiczno-prywatne w żadnym wypadku nie stanowi jedynego ani *a priori* lepszego rozwiązania i należy je uwzględniać tylko w sytuacji, gdy można wykazać, że jego zastosowanie przyniesie wartość dodaną w porównaniu z innymi podejściami, gdy może być skutecznie wdrażane oraz jeżeli wszystkie strony mogą osiągnąć swoje cele. **W Polsce ustawa o partnerstwie publiczno-prywatnym (art. 3) rozstrzyga, że może ono być sposobem realizacji inwestycji, jeśli przynosi to korzyści dla interesu publicznego przewyższające korzyści realizacji tego przedsięwzięcia w inny sposób**.

Prezentowane ujęcie zagadnienia korzyści ma w szczególności następujące konsekwencje (Amerykańska Izba Handlowa w Polsce 2002):

- kryterium korzyści dla interesu publicznego decyduje o zakwalifikowaniu poszczególnych inwestycji publicznych do realizacji w trybie partnerstwa publiczno-prywatnego,
- każdy przewidziany do realizacji projekt musi być poddany gruntownej ocenie pod kątem możliwych korzyści dla interesu publicznego,
- do oceny poszczególnych projektów niezbędne jest zastosowanie odpowiednich narzędzi, umożliwiających pomiar i porównanie korzyści, które można uzyskać, stosując różne metody realizacji danego przedsięwzięcia.

³ Więcej na ten temat w dalszej części artykułu.

⁴ Zagadnienie *value for money* stanowi treść ostatniego rozdziału artykułu.

Przed podjęciem decyzji o zastosowaniu partnerstwa publiczno-prywatnego, w trosce o interes publiczny, podmiot publiczny musi więc przeprowadzić stosowne analizy.

W partnerstwie publiczno-prywatnym dla uzyskania oczekiwanych korzyści decydujące znaczenie ma podział zadań i ryzyka pomiędzy sektor publiczny a prywatnych inwestorów. Powinien on umożliwić przyjęcie przez każdą ze stron tych rodzajów ryzyka, do których kontrolowania jest ona najlepiej przygotowana ze względu na posiadaną wiedzę, umiejętności i doświadczenie. Szczegóły podziału zadań i ryzyka w odniesieniu do każdego przedsięwzięcia zawiera umowa między podmiotami współpracującymi na zasadach partnerstwa publiczno-prywatnego, stanowiąca podstawę tej współpracy.

Partnerstwo publiczno-prywatne nie zwalnia sektora publicznego z odpowiedzialności za wykonanie danego projektu. Pozwala natomiast na przejęcie przez prywatnych partnerów mniejszej bądź większej części ryzyka projektu. Jest to zasadne jedynie wtedy, gdy inwestorzy prywatni są w stanie zarządzać określonym rodzajem ryzyka efektywniej niż sektor publiczny, czyli ponosząc mniejsze koszty.

3.2. Źródła korzyści wynikających z zastosowania partnerstwa publiczno-prywatnego

Jak już wspomniano, partnerstwo publiczno-prywatne umożliwia finansowanie zadań publicznych bez angażowania środków publicznych lub wykorzystanie ich w mniejszym stopniu niż w przypadku zastosowania metod tradycyjnych. Niewątpliwie przyczyniło się to do upowszechnienia partnerstwa publiczno-prywatnego w wielu krajach. **Dzięki zaangażowaniu prywatnych inwestorów w realizację zadań publicznych można uzyskać także inne korzyści, wynikające – najogólniej mówiąc – z właściwego im sposobu prowadzenia działalności gospodarczej. Wiedza, umiejętności i doświadczenie prywatnych przedsiębiorców, udostępniane przez nich w ramach partnerstwa publiczno-prywatnego, sprzyjają optymalizacji efektywności realizowanych inwestycji w całym cyklu ich życia** (ang. *whole life cycle approach*). W szczególności istotne wydaje się to, że (Amerykańska Izba Handlowa w Polsce 2002):

- prywatni inwestorzy mają dostęp do innowacyjnych technologii, których zastosowanie ułatwia rozbudowę nowoczesnej infrastruktury i poprawę jakości usług publicznych,
- zaangażowanie prywatnych inwestorów zwiększa prawdopodobieństwo przeprowadzenia inwestycji zgodnie z harmonogramem, bez opóźnień typowych dla podmiotów publicznych, mających mniejsze doświadczenie jako przedsiębiorcy,
- prywatni wykonawcy, dążący do osiągnięcia możliwie wysokiej stopy zwrotu, o wiele bardziej przestrze-

gają dyscypliny finansowej realizowanych projektów niż podmioty publiczne, które nie są przyzwyczajone do zachowań rynkowych.

Dostęp do zaawansowanych technologii

Pozyskanie i wykorzystanie zaawansowanych technologii, udostępnianych przez prywatnych inwestorów w ramach współpracy z sektorem publicznym umożliwia rozbudowę i rozwój infrastruktury na światowym poziomie. Dzięki temu jakość usług publicznych świadczonych za pomocą nowoczesnej infrastruktury może być odpowiednio wysoka.

Transfer zaawansowanych technologii do gospodarki publicznej realizowany w ramach partnerstwa publiczno-prywatnego ułatwia zaspokojenie rosnących oczekiwań społecznych co do poziomu usług oferowanych przez administrację publiczną. Sprzyja też budowie potencjału technologicznego państwa i zwiększeniu jego technologicznej samowystarczalności. Co ważne – nie wymaga to podejmowania przez państwo szczególnych wysiłków na rzecz uzyskania dostępu do innowacyjnych technologii ani angażowania dodatkowych, zawsze ograniczonych, środków publicznych. Korzyści, które można dzięki temu uzyskać, trudno przecenić. *Know-how* stosowane w partnerstwie publiczno-prywatnym to sposób na zmniejszenie kosztów realizacji projektów w całym cyklu ich życia.

Wykonanie zadania publicznego zgodnie z obowiązującym harmonogramem

Z doświadczenia państw stosujących partnerstwo publiczno-prywatne wynika, że zdecydowana większość projektów partnerstwa jest realizowana w terminie. Tymczasem inwestycje sektora publicznego realizowane metodami tradycyjnymi rzadko przebiegają bez opóźnień. Różnica ta wynika przede wszystkim z tego, że dla inwestorów prywatnych są to projekty komercyjne, które – jak wszystkie inne – mają zapewnić uzyskanie satysfakcjonującej stopy zwrotu. To, że celem przedsięwzięć podejmowanych w ramach partnerstwa publiczno-prywatnego jest realizacja zadań publicznych, nie zwalnia ich wykonawców z działania typowego dla projektów biznesowych (Stech 2008). **Opóźnienia, tak charakterystyczne dla działalności gospodarczej prowadzonej przez podmioty publiczne w sposób tradycyjny, powodujące marnotrawstwo środków publicznych, a czasami wręcz niewyobrażalne straty, w projektach partnerstwa publiczno-prywatnego nie mogą mieć miejsca.** Uczestnictwo podmiotów prywatnych istotnie zwiększa zatem prawdopodobieństwo wykonania poszczególnych zadań publicznych zgodnie z harmonogramem i tym samym uniknięcia kosztów, które powstają w przypadku przekroczenia terminów. Efektywność zadania publicznego wykonanego na czas przewyższa efektywność

przedsięwzięcia, w którego realizacji występują opóźnienia, i jest źródłem oczekiwanych korzyści.

Wykonanie zadania publicznego zgodnie z obowiązującym planem finansowym

Nieprzestrzeganie dyscypliny finansowej przy wykonywaniu zadań publicznych jest powszechnym zjawiskiem w przypadku ich realizacji przez sektor publiczny w sposób tradycyjny. Powoduje to przekraczanie budżetów poszczególnych inwestycji, generuje koszty, które muszą być pokryte z publicznych środków finansowych i które często są tak wysokie, że uniemożliwiają całe przedsięwzięcie. Dla interesu publicznego zdecydowanie korzystniejsze jest zatem wykonanie zadania publicznego w ramach obowiązującego budżetu. Z praktyki wynika jednak, że wśród tak zrealizowanych inwestycji zdecydowanie przeważają projekty prowadzone na zasadach partnerstwa publiczno-prywatnego, a więc z udziałem inwestorów prywatnych. Dzieje się tak m.in. dlatego, że inwestorzy prywatni, dążąc do maksymalizacji wyników, czują się zobowiązani planem finansowym, natomiast zadania publiczne realizowane według metod tradycyjnych rzadko mieszczą się w wyznaczonych granicach kosztów. Przestrzeganie budżetu zapewnia osiągnięcie większej efektywności niż w przypadku inwestycji, których budżety zostały przekroczone. Wydaje się więc oczywiste, że partnerstwo publiczno-prywatne znacznie zwiększa prawdopodobieństwo utrzymania dyscypliny finansowej i tym samym odniesienia korzyści, których uzyskanie nie jest możliwe w przypadku realizacji tego zadania przez sektor publiczny we własnym zakresie. Prawdopodobieństwo to wydaje się jeszcze większe, jeśli partnerzy prywatni uczestniczący w projekcie przejmują obowiązek jego sfinansowania. Zwalnia to sektor publiczny z wydatkowania środków publicznych na realizację przedsięwzięcia. Ponadto **inwestorzy prywatni efektywniej zarządzają kapitałem pozyskanym na realizację projektu, podczas gdy podmioty publiczne niestety zbyt często mają problemy z podporządkowaniem prowadzonej działalności rachunkowi ekonomicznemu.** Wprawdzie kapitał pozyskiwany przez inwestorów prywatnych (na przykład kredyt bankowy) z reguły jest droższy niż kapitał udostępniony sektorowi publicznemu (na przykład emisja obligacji skarbowych), jednak rekompensują to doświadczenie i umiejętności podmiotów prywatnych oraz dostępne im nowoczesne technologie.

Racjonalne postępowanie inwestorów prywatnych zaangażowanych we współpracę z sektorem publicznym w ramach partnerstwa publiczno-prywatnego, wsparte wiedzą i doświadczeniem sektora publicznego, z reguły prowadzi do uzyskania korzyści, których podmioty publiczne samodzielnie nie są w sta-

nie wygenerować. Z dotychczasowych doświadczeń w tym zakresie, np. Wielkiej Brytanii, wynika, że korzyści te są osiągane w przypadku zastosowania różnych modeli partnerstwa publiczno-prywatnego.

4. Ryzyko w projektach partnerstwa publiczno-prywatnego

4.1. Cel podziału ryzyka w partnerstwie publiczno-prywatnym

Dla skali korzyści możliwych do uzyskania w wyniku realizacji zadań publicznych w ramach partnerstwa publiczno-prywatnego, podstawowe znaczenie – jak już wspomniano – ma podział właściwego im ryzyka między strony realizujące dane przedsięwzięcie, związany z przeniesieniem niektórych rodzajów ryzyka z sektora publicznego na sektor prywatny. **Za optymalny można uznać taki podział ryzyka projektu, który prowadzi do maksymalizacji korzyści wynikających z jego realizacji** (HM Treasury 2003).

Możliwość przekazania partnerom prywatnym części lub całości ryzyka danej inwestycji, z jednoczesnym zachowaniem przez podmiot publiczny odpowiedzialności za pomyślną realizację wyznaczonego celu, jest uważana za jedną z podstawowych zalet partnerstwa publiczno-prywatnego. **Transfer ryzyka projektu pozwala zmienić rolę sektora publicznego z wykonawcy projektu na jego zarządcę. Nie jest on jednak celem samym w sobie. Jest zasadny, jeśli prowadzi do zwiększenia korzyści wynikających z realizacji danego przedsięwzięcia.**

Komisja Europejska (2003) definiuje cele przeniesienia ryzyka w następujący sposób:

- zmniejszenie kosztów projektu w długim okresie dzięki przeniesieniu ryzyka na stronę, która potrafi nim efektywniej zarządzać,
- zachęcenie wykonawcy do zrealizowania inwestycji na czas, w określonym standardzie i w ramach budżetu,
- poprawa jakości usług i zwiększenie dochodów poprzez bardziej rentowne działanie,
- zapewnienie bardziej spójnego i przewidywalnego profilu wydatków.

Zasada, którą kierują się strony umowy o partnerstwie publiczno-prywatnym przy podziale ryzyka projektu, wydaje się czytelna – każda strona przejmuje ten rodzaj ryzyka, którym jest w stanie lepiej zarządzać, co prowadzi do minimalizacji kosztów. Punktem wyjścia do podziału ryzyka jest właściwa identyfikacja rodzajów ryzyka charakteryzujących dany projekt, która ułatwia optymalny podział.

4.2. Rodzaje ryzyka w partnerstwie publiczno-prywatnym

Rodzaje ryzyka związane z realizacją przedsięwzięć w ramach partnerstwa publiczno-prywatnego są przed-

miotem różnych opracowań, o mniej lub bardziej formalnym charakterze.

Komisja Europejska (2003) zdefiniowała następujące rodzaje ryzyka związane z realizacją przedsięwzięć w ramach partnerstwa publiczno-prywatnego:

- ryzyko związane z dochodem,
- ryzyko wyboru partnera z sektora prywatnego,
- ryzyko związane z pracami budowlanymi,
- ryzyko walutowe,
- ryzyko regulacyjne (umowne),
- ryzyko polityczne,
- ryzyko środowiskowe (archeologiczne),
- ryzyko ukrytych wad,
- ryzyko związane ze społeczną akceptacją,
- ryzyko utraty trwałości,
- ryzyko akceptacji społecznej.

Jakiś czas później Eurostat, podejmując inicjatywę uporządkowania złożonej i nieco pogmatwanej kwestii rodzajów ryzyka, w tym stosowanego w tym zakresie nazewnictwa, będącego często źródłem dodatkowego zamieszania, zaproponował wyodrębnienie trzech podstawowych rodzajów ryzyka w odniesieniu do partnerstwa publiczno-prywatnego (Eurostat 2004):

– **ryzyko związane z budową** (ang. *construction risk*), obejmujące na przykład terminy związane z wykonaniem inwestycji, uzgodnione standardy wykonania prac, nieplanowane koszty, wady techniczne czy wpływ czynników zewnętrznych na przebieg realizacji przedsięwzięcia,

– **ryzyko związane z dostępnością** (ang. *availability risk*), obejmujące możliwość wystąpienia zakłóceń w udostępnianiu usług ostatecznym użytkownikom w sposób zgodny z warunkami umowy,

– **ryzyko związane z popytem** (ang. *demand risk*), obejmujące możliwość, że popyt na usługę oferowaną w ramach partnerstwa publiczno-prywatnego będzie mniejszy lub większy niż popyt przewidziany w umowie.

W Polsce rodzaje ryzyka związanego z partnerstwem publiczno-prywatnym określa Rozporządzenie Ministra Gospodarki w sprawie ryzyk związanych z realizacją przedsięwzięć w ramach partnerstwa publiczno-prywatnego. Polski ustawodawca, definiując rodzaje ryzyka w przepisach wykonawczych do ustawy o partnerstwie publiczno-prywatnym, zaproponował rozwiązanie znacznie bardziej szczegółowe od opracowanego przez Komisję Europejską. Wyodrębnił następujące rodzaje ryzyka⁵:

– **ryzyko związane z budową** – ryzyko powodujące zmiany kosztów i terminów związanych z realizacją budowy, przebudowy lub rozbudowy w ramach przedsięwzięcia, jak również z funkcjonowaniem udostępnionych już składników majątkowych,

– **ryzyko związane z dostępnością** – ryzyko wpływające na sposób, jakość lub ilość dostarczanych w ramach realizacji umowy o partnerstwie publiczno-prywatnym usług,

– **ryzyko związane z popytem** – ryzyko powodujące zmianę popytu na określone usługi,

– **ryzyko związane z przygotowaniem przedsięwzięcia** – ryzyko wpływające na koszt i czas trwania procesu przetargowego,

– **ryzyko rynkowe związane z dostępnością nakładów na realizację przedsięwzięcia** – ryzyko wpływające na koszt, ilość, jakość i terminy dostarczania nakładów niezbędnych do realizacji przedsięwzięcia,

– **ryzyko o charakterze politycznym** – ryzyko wystąpienia zmian w sferze polityki, której kierunki związane są z rozwojem przedsięwzięć realizowanych w ramach partnerstwa publiczno-prywatnego,

– **ryzyko o charakterze legislacyjnym** – ryzyko wystąpienia zmian w przepisach prawnych, mających wpływ na realizację przedsięwzięć w ramach partnerstwa publiczno-prywatnego,

– **ryzyko makroekonomiczne** – ryzyko wpływające na sytuację ekonomiczną,

– **ryzyko regulacyjne** – ryzyko wystąpienia zmian w regulacjach dotyczących systemów opłat w ramach danej dziedziny usług użyteczności publicznej, które mają wpływ na koszty realizacji przedsięwzięcia lub w wyniku których zmniejszenie ulegnie zakres praw i obowiązków stron w ramach przedsięwzięcia,

– **ryzyko związane z przychodami z przedsięwzięcia** – ryzyko mające wpływ na poziom przychodów uzyskiwanych w ramach realizacji przedsięwzięcia,

– **ryzyko związane z wystąpieniem siły wyższej,**

– **ryzyko związane z rozstrzygnięciem sporów** – ryzyko, którego wystąpienie wpływa na sposób i efektywność rozstrzygnięcia sporu powstałego na tle realizacji umowy o partnerstwie publiczno-prywatnym,

– **ryzyko związane ze stanem środowiska naturalnego** – ryzyko powodujące powstanie obowiązku podjęcia działań mających na celu poprawę środowiska naturalnego przed rozpoczęciem realizacji przedsięwzięcia lub ryzyko pogorszenia się stanu środowiska naturalnego w wyniku realizacji przedsięwzięcia,

– **ryzyko związane z lokalizacją przedsięwzięcia** – ryzyko wpływające na dostępność terenu przeznaczonego do realizacji przedsięwzięcia,

– **ryzyko związane z przekazywaniem składników majątkowych** – ryzyko wpływające na warunki i terminy przekazywania składników majątkowych w ramach realizacji przedsięwzięcia,

– **ryzyko związane z końcową wartością składników majątkowych** – ryzyko wartości materialnej składników majątkowych na dzień zakończenia realizacji umowy o partnerstwie publiczno-prywatnym,

– **ryzyko związane z brakiem społecznej akceptacji** – ryzyko protestów i sprzeciwów społeczności lo-

⁵ Rozporządzenie Ministra Gospodarki z dnia 21 czerwca 2006 r. w sprawie ryzyk związanych z realizacją przedsięwzięć w ramach partnerstwa publiczno-prywatnego, Dz.U. nr 125, poz. 866.

kalnych, w szczególności podczas wdrażania i realizacji projektów infrastrukturalnych w ramach partnerstwa publiczno-prywatnego.

W polskim ustawodawstwie lista rodzajów ryzyka jest długa w porównaniu z rozstrzygnięciem Komisji Europejskiej. Trzy rodzaje ryzyka wyodrębnione przez Eurostat (ryzyko budowy, dostępności i popytu) otwierają katalog siedemnastu pozycji z rozporządzenia ministra gospodarki. Wszystkich wymienionych w rozporządzeniu rodzajów ryzyka jest więcej, niż przedstawiono powyżej, ponieważ większość pozycji ustawodawca rozszerzył o przykłady konkretnych rodzajów ryzyka⁶.

4.3. Podział ryzyka w partnerstwie publiczno-prywatnym

W przypadku każdego zadania publicznego realizowanego w ramach partnerstwa publiczno-prywatnego podział ryzyka między sektor publiczny a podmioty prywatne powinien uwzględniać właściwości projektu. Punktem wyjścia podziału ryzyka jest podział między strony zadań i obowiązków wynikających z podejmowanego przedsięwzięcia. Od sposobu podziału zadań i – w konsekwencji – ryzyka zależą korzyści możliwe do uzyskania w wyniku realizacji danego projektu. Z tego względu **alokację ryzyka można uznać za podstawową kwestię w każdym przedsięwzięciu wykonywanym w ramach partnerstwa publiczno-prywatnego, decydującą o jego opłacalności**. Dokonując podziału ryzyka, strony powinny uwzględniać swe przygotowanie do zarządzania danym rodzajem ryzyka w celu osiągnięcia jak największych korzyści.

Przeniesienie ryzyka na prywatnych inwestorów nie jest celem samym w sobie. Jest zasadne, jeśli prowadzi do zwiększenia możliwości do uzyskania korzyści. W przeciwnym przypadku pozostawienie ryzyka po stronie sektora publicznego powinno być efektywniej-

sze⁷. Sektor prywatny przejmuje więc te rodzaje ryzyka i w takim stopniu, by móc zarządzać nimi efektywniej od sektora publicznego. Ryzyko przejmowane przez sektor prywatny często jest rozłożone na różne podmioty, uczestniczące w realizacji przedsięwzięcia w ramach partnerstwa publiczno-prywatnego, zgodnie z ich kompetencjami. **Przejęcie ryzyka przez prywatnych inwestorów zwiększa ich motywację do uzyskania jak najlepszych wyników w efekcie realizacji projektu, co prowadzi do maksymalizacji korzyści dla interesu publicznego.**

Czasem nie jest łatwo jednoznacznie rozstrzygnąć, która ze stron umowy o partnerstwie publiczno-prywatnym odpowiada za dany rodzaj ryzyka. Wówczas z reguły przyjmuje się, że dany rodzaj ryzyka pozostaje pod kontrolą tej strony, na którą przypada jego większa część (np. Eurostat 2004). W niektórych przypadkach najlepszym rozwiązaniem może okazać się wspólne zarządzanie danym rodzajem ryzyka przez oba sektory. Tabela 1 zawiera przykładowy podział ryzyka w między sektor publiczny a partnerów prywatnych w projekcie partnerstwa publiczno-prywatnego.

Podział ryzyka w partnerstwie publiczno-prywatnym a poziom długu publicznego i deficytu sektora finansów publicznych

Podział ryzyka między sektor publiczny a partnerów prywatnych, wspólnie wykonujących zadanie publiczne na podstawie umowy o partnerstwie publiczno-prywatnym, skutkuje powstaniem określonych zobowiązań. **Dla skali korzyści sektora publicznego, możliwych do uzyskania w wyniku realizacji inwestycji, podstawowe znaczenie ma wpływ zobowiązań, wynikających z umowy o partnerstwie publiczno-prywatnym, na poziom długu publicznego i deficytu jednostek sektora finansów publicznych, zależny właśnie od dokonanego podziału ryzyka między strony umowy.**

⁶ Na przykład ryzyko makroekonomiczne, zgodnie z treścią rozporządzenia, to w szczególności ryzyko: inflacji, zmian wysokości stóp procentowych, kursowe, zmian demograficznych, związane z tempem rozwoju gospodarczego. Więcej szczegółów dotyczących rodzajów ryzyka zawiera przytaczane wcześniej rozporządzenie ministra gospodarki.

⁷ Rodzaje ryzyka, których transfer na podmioty prywatne nie będzie prowadził do dodatkowych korzyści to np. ryzyko polityczne, ryzyko inflacji czy ryzyko zmiany potrzeb społecznych w trakcie obowiązywania umowy o partnerstwie publiczno-prywatnym.

Tabela 1. Przykładowy podział ryzyka w projekcie partnerstwa publiczno-prywatnego

Rodzaj ryzyka	Partnerzy prywatni	Sektor publiczny
Ryzyko przekroczenia terminu prac budowlanych	•	
Ryzyko wzrostu kosztów	•	
Ryzyko wystąpienia zmian technologicznych	•	
Ryzyko o charakterze legislacyjnym		•
Ryzyko kursowe	•	
Ryzyko polityczne		•
Ryzyko związane ze stanem środowiska		•
Ryzyko siły wyższej	•	•

Źródło: opracowanie własne na podstawie Rozporządzenia Ministra Gospodarki z dnia 21 czerwca 2006 r. w sprawie ryzyk związanych z realizacją przedsięwzięć w ramach partnerstwa publiczno-prywatnego, Dz.U. nr 125, poz. 866.

W szczególności to, czy zobowiązania wynikające z umowy o partnerstwie publiczno-prywatnym spowodują zwiększenie długu i deficytu publicznego, czy nie, zależy od sposobu podziału ryzyka związanego z (Eurostat 2004):

- budową,
- dostępnością,
- popytem.

W polskim ustawodawstwie zasady podziału tych trzech rodzajów ryzyka zostały określone w Rozporządzeniu Ministra Gospodarki w sprawie ryzyk wynikających z realizacji przedsięwzięć w ramach partnerstwa publiczno-prywatnego, w ślad za decyzją Eurostatu w sprawie traktowania zobowiązań wynikających z realizacji przedsięwzięć w ramach partnerstwa publiczno-prywatnego i ich wpływu na dług i deficyt publiczny. W rozporządzeniu znajduje się następujący zapis (paragraf 4):

„Zobowiązania wynikające z umowy o partnerstwie publiczno-prywatnym mają wpływ na poziom długu publicznego oraz deficytu sektora finansów publicznych w przypadku, gdy podmiot publiczny:

- 1) ponosi większość ryzyk związanych z budową albo
- 2) ponosi równocześnie większość ryzyk związanych z dostępnością i ryzyk związanych z popytem.

Podstawowym narzędziem oceny podziału ryzyka poszczególnych projektów i ewentualnego wpływu tego podziału na poziom długu publicznego i deficytu finansów publicznych jest analiza jakościowa ryzyka. Obejmuje ona określenie: podziału ryzyka projektu pomiędzy strony umowy partnerstwa publiczno-prywatnego, prawdopodobieństwo wystąpienia odrębnych rodzajów ryzyka oraz ich wpływu na realizowane przedsięwzięcie. W niektórych przypadkach celowe jest uzupełnienie analizy jakościowej ryzyka jego analizą ilościową. Polega ona na liczbowym wyrażeniu wszystkich zidentyfikowanych rodzajów ryzyk przedsięwzięcia na podstawie prawdopodobieństwa ich wystąpienia oraz ich przewidywanego wpływu na przedsięwzięcie.

Niekiedy ze względu na mało przejrzysty podział ryzyka nie jest możliwe jednoznaczne wskazanie, czy większością poszczególnych rodzajów ryzyka zarządza sektor publiczny, czy sektor prywatny. Trudno przez to określić wpływ zobowiązań wynikających z umowy o partnerstwie publiczno-prywatnym na stan finansów publicznych. Po wykorzystaniu dodatkowych narzędzi oceny ryzyka często okazuje się, że zobowiązania te zwiększą dług i deficyt publiczny.

Skala korzyści sektora publicznego możliwych do uzyskania w wyniku realizacji danego przedsięwzięcia w ramach partnerstwa publiczno-prywatnego zależy głównie, jak już wcześniej wspomniano, od podziału ryzyka projektu między sektor publiczny i partnerów prywatnych oraz wynikających z tego podziału skutków dla sektora finansów publicznych. Określenie skali korzyści pomaga zakwalifikować dany projekt do realizacji według formuły partnerstwa publiczno-prywatnego lub w inny sposób.

5. Value for money jako miara korzyści dla interesu publicznego wynikających z realizacji przedsięwzięcia w ramach partnerstwa publiczno-prywatnego

Do szacowania korzyści z wykonania zadań publicznych w ramach partnerstwa publiczno-prywatnego stosowane są złożone techniki. **Powszechnie stosowaną miarą korzyści jest value for money⁸.**

Wartościowe ujęcie korzyści możliwych do uzyskania w wyniku realizacji przedsięwzięć w różny sposób jest niezbędne do porównania tych korzyści i wyboru metody pozwalającej uzyskać największą ich wartość. Zgodnie ze wspomnianym już podstawowym kryterium określania sposobu realizacji zadań publicznych partnerstwo publiczno-prywatne powinno prowadzić do osiągnięcia korzyści większych niż możliwe do uzyskania w inny sposób.

W ocenie korzyści pomocne jest zdefiniowanie tego pojęcia, dzięki czemu łatwiejsze staje się rozpoznanie korzyści, a następnie ich wartościowe ujęcie. Różne państwa stosują odmienne definicje korzyści (opłacalność, *value for money*) z partnerstwa publiczno-prywatnego, jednak dla wielu wspólne jest zaakcentowanie relatywnej przewagi partnerstwa publiczno-prywatnego nad tradycyjnymi metodami realizacji przedsięwzięcia (np. HM Treasury 2004). Polski ustawodawca uniknął zdefiniowania pojęcia korzyści wprost, ograniczając się do podania przykładów sytuacji, w których korzyści występują: oszczędność w wydatkach pomiotu publicznego, podniesienie standardu świadczonych usług lub obniżenie uciążliwości dla otoczenia.

Ocenę opłacalności projektów partnerstwa publiczno-prywatnego należy prowadzić na kolejnych etapach przygotowania inwestycji do realizacji. Można w tym celu wykorzystać np. następujące narzędzia oceny ilościowej:

- PPP Potencial Scan (PPS),
- Public Private Comparator (PPC),
- Public Sector Comparator (PSC),

które w różnych państwach mogą mieć nieco odmienną treść i odgrywać różną rolę w ocenie opłacalności partnerstwa publiczno-prywatnego (ECORYS 2004; Glinka 2006).

PPP Potencial Scan (PPS)

PPS jest narzędziem wstępnej oceny projektu pod kątem jego zdolności do wygenerowania korzyści (opłacalności, *value for money*) w przypadku zastosowania partnerstwa publiczno-prywatnego. Służy do wyodrębnienia dwóch kategorii projektów:

- projektów, które ze względu na możliwe do uzyskania korzyści powinny zostać zrealizowane w formule partnerstwa publiczno-prywatnego,

⁸ Termin *value for money* właściwie nie ma jednego odpowiednika polskiego. Często jest tłumaczony jako „korzyści” (tego terminu używa np. polski ustawodawca) lub „opłacalność”. Często jest też używany w oryginalnym brzmieniu.

– projektów, w odniesieniu do których lepszym, bo gwarantującym większe korzyści, jest tradycyjny sposób ich realizacji.

Public Private Comparator (PPC)

W odniesieniu do przedsięwzięć, o których już wiadomo, że będą realizowane na zasadach partnerstwa publiczno-prywatnego wybiera się model partnerstwa, który da największe korzyści. W tym celu należy:

- po pierwsze, dokonać wyboru ewentualnego sposobu wykonania zadania publicznego w sposób tradycyjny (model publiczny) oraz określić korzyści, które wystąpiłyby w przypadku jego zastosowania,
- po drugie, określić warianty partnerstwa publiczno-prywatnego możliwe do zastosowania oraz korzyści, które wystąpiłyby w przypadku zastosowania każdego z nich,
- po trzecie, porównać korzyści możliwe do uzyskania dla tych wariantów z korzyściami, które powstałyby, gdyby zadanie zostało wykonane zgodnie z modelem publicznym,
- po czwarte, **wybrać najlepszy wariant partnerstwa publiczno-prywatnego, czyli taki, dla którego różnica między korzyściami z jego zastosowania a korzyściami z modelu publicznego będzie największa (największa wartość *value for money*).**

W przypadku uzyskania przez projekt ujemnej wartości PPC, nawet mimo wcześniej wykazanego dodatniego PPS, wskazane jest odstąpienie od partnerstwa publiczno-prywatnego i rozważenie realizacji danego przedsięwzięcia z wykorzystaniem modelu publicznego.

Public Sector Comparator (PSC)

Trzeci, ostatni etap oceny opłacalności to wybór najkorzystniejszej oferty sektora prywatnego, czyli dającej największą wartość *value for money*, dokonany zgodnie z kryterium PSC.

Public Sector Comparator odpowiada hipotetycznym kosztom wykonania zadania publicznego przez sektor publiczny bez współpracy z inwestorami prywatnymi. Jest zatem wartością referencyjną kosztów realizacji projektu i umożliwia porównanie opłacalności rozwiązań w ramach wybranego wariantu partnerstwa publiczno-prywatnego z korzyściami wynikającymi z zastosowania metody tradycyjnej. Jeśli koszty oferty prywatnego inwestora są niższe od wartości PSC, to jest

to oferta zdolna do wykreowania *value for money* (zob. schemat 1). Porównanie PSC z kosztami wszystkich zgłoszonych ofert prywatnych inwestorów umożliwia wybór rozwiązania zapewniającego największą *value for money*, czyli najlepszego (Glinka 2006).

Stosowanie kryterium *value for money* do oceny opłacalności partnerstwa publiczno-prywatnego w wykonywaniu zadań publicznych nie jest proste m.in. dlatego, że:

- wymaga długiego okresu – może trwać nawet dłużej niż rok (zob. tabela 2),
- wymaga dużych umiejętności i odpowiedniego doświadczenia osób prowadzących ocenę, której tryb i procedury są złożone i kosztowne,
- **wyniki przeprowadzonej analizy nie są obiektywne** i w znacznym stopniu zależą od jakości dostępnych danych, wykorzystanych do oceny.

W wielu przypadkach konieczne jest uzupełnienie oceny ilościowej partnerstwa publiczno-prywatnego analizą jakościową. Uwzględnić się w niej czynniki mające wpływ na realizację przedsięwzięcia i niepoddające się ocenie ilościowej, na przykład kwestie społeczne czy jakość informacji stanowiących podstawę analizy ilościowej projektu. Ocena jakościowa jest tym bardziej pożądana, im mniejsza jest wykazana przewaga partnerstwa publiczno-prywatnego nad metodą tradycyjną, mierzona *value for money*.

Schemat 1. *Value for money projektu partnerstwa publiczno-prywatnego*

PSC > koszty realizacji projektu w ramach partnerstwa publiczno-prawnego → *value for money*

Źródło: opracowanie własne na podstawie: ECORYS (2004, część II).

6. Podsumowanie

Partnerstwo publiczno-prywatne umożliwia udział inwestorów prywatnych w wykonywaniu zadań publicznych, a w szczególności w ponoszeniu związanego z tym ryzyka. Możliwość przeniesienia ryzyka na partnerów prywatnych, w połączeniu z ich znacznie lepszym przygotowaniem do prowadzenia działalno-

Tabela 2. Orientacyjny koszt trwania oceny w odniesieniu do poszczególnych narzędzi

Narzędzie oceny	Czas trwania oceny	Etap przygotowania projektu do realizacji
PPS	0,5–1,5 miesiąca	(bardzo) wstępny
PPC	2–4 miesiące	przygotowanie przetargu
PSC	4–8 miesięcy	trwanie przetargu

Źródło: ECORYS (2004, część V, s. 20–22).

ści gospodarczej jest podstawowym źródłem korzyści dla interesu publicznego, których sektor publiczny nie jest w stanie uzyskać, wykonując zadania publiczne we własnym zakresie. Zastosowanie partnerstwa publiczno-prywatnego z reguły prowadzi do zwiększenia dostępności usług publicznych, jak również poprawy ich jakości, co potwierdzają doświadczenia wielu krajów.

Ocena zasadności realizacji przedsięwzięć w trybie partnerstwa publiczno-prywatnego, zgodnie z podstawowym kryterium *value for money*, zabiera dużo czasu i pochłania znaczne środki, co w przypadku wielu projektów skutecznie zniechęca do jej podjęcia. Z drugiej strony jest to sposób na znalezienie tych przedsięwzięć, w odniesieniu do których zastosowanie określonych modeli partnerstwa publiczno-prywatnego będzie prowadzić do uzyskania *value for money*. Dzięki temu **społeczeństwo zyskuje dostęp do nowoczesnej infrastruktury i usług o wysokiej jakości**⁹. Zdrowie, oświata, kultura, rekreacja, transport to obszary, gdzie wyniki zaangażowania prywatnych inwestorów w wykonywanie zadań publicznych są widoczne. **Dlatego warto zabiegać o rozwój partnerstwa publiczno-prywatnego.**

W Polsce stan infrastruktury i zaniedbania w zakresie usług publicznych wymagają podjęcia jak naj-

szybszych działań naprawczych. Wiele wskazuje na to, że do odrabiania zaległości można wykorzystać partnerstwo publiczno-prywatne. Do tego konieczne jest jednak przełamanie istniejących barier, a to z kolei wymaga aktywnej i konsekwentnej postawy rządu. Bez podjęcia przez rząd określonych inicjatyw rozwój partnerstwa publiczno-prywatnego nie będzie możliwy.

Warunki rozwoju partnerstwa publiczno-prywatnego w Polsce nie wydają się gorsze niż w innych państwach, którym udało się pokonać trudności i które w mniejszym lub większym stopniu korzystają z wyników współpracy sektora publicznego z inwestorami prywatnymi. W Polsce kolejne ekipy rządzące sprawiają wrażenie niezainteresowanych wprowadzeniem partnerstwa publiczno-prywatnego. Jeśli nawet były podejmowane jakieś starania, to okazywały się całkowicie nieskuteczne, czego koronnym przykładem jest choćby ustawa o partnerstwie publiczno-prywatnym, która nie jest stosowana.

Przykład wielu państw wskazuje, że warto ponieść trud i wysokie koszty upowszechnienia partnerstwa publiczno-prywatnego. Jego rozwój zwiększa bowiem skuteczność wysiłków rządu na rzecz poprawy jakości życia społeczeństwa. Przedmiotem dalszych badań autorki będzie stan partnerstwa publiczno-prywatnego w Polsce oraz perspektywy jego rozwoju, jak również doświadczenia innych państw w zakresie wykorzystania partnerstwa do realizacji zadań publicznych.

⁹ Usługi publiczne oferowane w ramach partnerstwa publiczno-prywatnego często są nie tylko lepszej jakości, niż usługi świadczone tradycyjną metodą, ale również od nich droższe. Ta potencjalnie wyższa cena jest argumentem chętnie wykorzystywanym przez przeciwników stosowania partnerstwa publiczno-prywatnego.

Bibliografia

- Allen G. (2003), *Private Finance Initiative (PFI)*, "Research Paper", No. 03/79, House of Commons, Economic Policy and Statistic Section, <http://www.parliament.uk/commons/lib/research/rp2003/rp03-079.pdf>
- Amerykańska Izba Handlowa w Polsce (2002), *Partnerstwo publiczno-prywatne jako metoda rozwoju infrastruktury w Polsce. Raport Amerykańskiej Izby Handlowej w Polsce*, Warszawa.
- Bates M. (1997), *Review of PFI (Private/Public Partnerships), Summary and Conclusions*, HM Treasury, London.
- British Columbia Ministry of Municipal Affairs (1999), *Public Private Partnership. A Guide for Local Government*, May, http://www.cserv.gov.bc.ca/lgd/policy_research/library/public_private_partnerships.pdf
- CBI (2007a), *Building on Success: The Way Forward for PFI*, Confederation for British Industry, <http://cbi.org.uk/pdf/healthreport0607.pdf>
- CBI (2007b), *Going Global: The World of Public-Private Partnerships*, Confederation of British Industry, <http://www.cbi.org.uk/pdf/goingglobal0707.pdf>
- Deloitte (2006a), *Closing the Infrastructure Gap. The role of Public-Private Partnerships. A Deloitte Research Study*, [http://www.deloitte.com/dtt/cda/doc/content/us_ps_ClosingInfrastructureGap2006\(1\).pdf](http://www.deloitte.com/dtt/cda/doc/content/us_ps_ClosingInfrastructureGap2006(1).pdf)
- Deloitte (2006b), *PROJECT 1p/05-3 Application of the PPP principle on the Economic and Social Cohesion Policy. Final Report for the Ministry for Regional Development of the Czech Republic*, May, <http://www.pppcentrum.cz/res/data/003/000488.pdf>
- Eaton D., Akbiyikli R. (2005), *Quantifying quality. A report on PFI and the delivery of public services*, Royal Institution for Chartered Surveyors, Coventry.
- ECORYS (2004), *Tworzenie podstaw dla PPP w Polsce*, ECORYS, Warszawa.
- European Commission (2004), *Resource Book on PPP Case Studies*, czerwiec, Directorate General for Regional Policy, Brussels.
- Eurostat (2004), *New decision of Eurostat on deficit and debt. Treatment of public-private partnerships*, 11 February, Luxembourg.
- Glinka J. (2006), *Przygotowanie partnerstwa publiczno-prywatnego*, prezentacja z konferencji „Zamówienia publiczne w projektach partnerstwa publiczno-prywatnego. Wybór partnera prywatnego”, MGG, 1 lutego, Warszawa.
- HM Treasury (2003), *PFI: Meeting the investment challenge*, July, London.
- HM Treasury (2004), *Value for Money Assessment Guidance*, August, London.
- HM Treasury (2006), *PFI: Strengthening long-term partnerships*, March, London.
- HM Treasury (2007), *Standardisation of PFI contracts*, March, HM Treasury, London.
- Komisja Europejska (2003), *Wytyczne dla udanego partnerstwa publiczno-prywatnego*, styczeń, Dyktoriat Generalny Polityki Regionalnej, http://www.ipppl.pl/instytut/temp/tlumaczenie_wytycznych_ke_dot_ppp.pdf
- Korbus B.P., Strawiński M. (2006), *Partnerstwo publiczno-prywatne. Nowa forma realizacji zadań publicznych*, Wydawnictwo Prawnicze LexisNexis, Warszawa.
- Moszoro M. (2005), *Partnerstwo publiczno-prywatne w monopolach naturalnych w sferze użyteczności publicznej*, Oficyna Wydawnicza SGH, Warszawa.
- NAO (1998a), *The PFI Contracts for Bridgend and Fazakerley Prisons*, National Audit Office, <http://www.nao.org.uk/pn/9798253.htm>
- NAO (1998b), *The Private Finance Initiative: The First Four Design, Build, Finance and Operate Road Contracts*, National Audit Office, <http://www.nao.org.uk/pn/9798476.htm>
- PPP Unit (2003), *Standardised PPP Provisions ("Standardisation")*, May, <http://www.ppp.gov.za/StandPPPProv.htm>
- State Government of Victoria Department of Treasury and Finance (2001), *Partnerships Victoria. Practitioners' Guide. Guidance Material*, March, <http://www.partnerships.vic.gov.au/>
- Stech G. (2008), *Jak nie należy budować autostrad w Polsce. 3,2 mld zł za opóźnienia*, „Businessman”, nr 2, s. 20–24.
- The Stationery Office (2000), *Public Private Partnerships. The Government's Approach*, London.
- UNIDO (2006), *BOT w projektach partnerstwa publiczno-prywatnego*, Difin, Warszawa.

Peter L. Bernstein, *Capital Ideas Evolving*

Review of the book by Peter L. Bernstein, *Capital Ideas Evolving*

John Wiley & Sons, Inc., Hoboken, New Jersey 2007

Adam Koronowski*

Capital Ideas Evolving jest nową książką Petera L. Bernsteina, finansisty od kilku dziesięcioleci zaangażowanego w doradztwo i zarządzanie finansowe oraz autora wcześniejszych książek: *Capital Ideas. The Improbable Origins of Modern Wall Street* (1992), *Against the Gods: The Remarkable Story of Risk* (1996), *The Power of Gold: The History of an Obsession* (2000)¹. Nowa książka już tytułem nawiązuje do *Capital Ideas*, w której Bernstein przedstawiał narodziny i wzrost uznania na rynkach finansowych koncepcji z zakresu teorii finansów, które obejmuje łącznym mianem *capital ideas*, a które w tej recenzji dalej nazywam klasyczną teorią finansów. Na ten zbiór składa się przede wszystkim teoria portfelowa Markowitza, model CAPM Sharpa, idea rynku efektywnego Fama wraz z racjonalnymi oczekiwaniami oraz model wyceny opcji Blacka-Scholesa-Mertona. W nowej książce Bernstein stawia sobie za cel ukazanie ewolucji akademickiej klasycznej teorii finansów dokonującej się poprzez jej praktyczne zastosowania na rynku finansowym. Siłą książki jest osobista znajomość łącząca Bernsteina ze wszystkim wymienionymi znakomitościami (w większości laureatami nagrody Nobla) i oparcie treści książki w znacznej mierze na rozmowach z wybitnymi teoretykami i praktykami rynku finansowego.

Trzeba jednak wyraźnie – znacznie mocniej, niż czyni to Bernstein – zaznaczyć, że praktyczne walory klasycznej teorii finansów są w świetle doświadczenia

i badań empirycznych wątpliwe lub bardzo skromne. Typowe dla postawy Bernsteina jest stwierdzenie, że „być może najważniejszą cechą tych idei jest ich przemożny wpływ na decyzje inwestycyjne, mimo że te teorie nie oparły się licznym testom empirycznym” (s. XVIII²). Bernstein wbrew faktom nie przestaje być zafascynowany klasyczną teorią i nie jest w tym odosobniony.

Liczne obserwacje zjawisk niezgodnych z rozumieniem racjonalności przyjętym przez teorię klasyczną oraz wyprowadzoną w jej ramach koncepcją rynku efektywnego stanowiły zaczyn, na którym wyrosła behawioralna teoria finansów³. Bernstein traktuje tę teorię jako opis anomalii, odstępstw od teorii klasycznej, które opisane i zrozumiane staną się przedmiotem korygujących działań przywracających rynkom efektywność. Owe anomalie miałyby być jedynie źródłem stóp zwrotu wyższych niż zgodne z teorią klasyczną, tzn. półmistycznej alfy, i prowadzić do pędu inwestorów eliminującego ową alfę. Bernstein ujmuje to następująco: „Każda alfa ścigana przez nienasycone rekiny musi mieć krótki żywot. Dlaczego rynki nie miałyby osiągnąć stanu, w którym rekiny pożarły wszystko w zasięgu wzroku? Wtedy otrzymalibyśmy rynki doskonale efektywne, gdzie ryzyko znajduje właściwe odzwierciedlenie w stopie zwrotu. Osiągnęlibyśmy (...) nirwanę (...)” (s. 244).

W sporze pomiędzy klasyczną a behawioralną teorią finansów Bernstein zajmuje pozycję stronniczego ar-

¹ Wszystkie trzy książki ukazały się w języku polskim wydane przez warszawskie wydawnictwo WIG-Press: pierwsza pt. *Intelektualna historia Wall Street* (1998), druga pt. *Przeciw bogom: niezwykłe dzieje ryzyka* (1997), trzecia pt. *Historia złota: dzieje obsesji* (2003)

² Wszystkie tłumaczenia zamieszczone w opracowaniu są autorstwa recenzenta.

³ Obszerne omówienie tej teorii polski Czytelnik znajdzie w Cieślak (2003).

bitra. Nie może zaprzeczyć spostrzeżeniom teorii behawioralnej, ale za wieczny i doskonały paradygmat uznaje teorię klasyczną. Taka postawa wyraża bardzo uproszczone postrzeganie dylematu, na którym koncentruje się, nie zawsze konsekwentnie, autor recenzowanej książki. Poniżej postaram się wyjaśnić, na czym to błędne uproszczenie polega.

Zacznijmy od stwierdzenia, że teoria portfelowa H. Markowitza, niejako otwierająca teorię klasyczną, ukazuje w ścisły i przejrzysty sposób korzyści z dywersyfikacji portfela i warunki ich maksymalizacji (konstrukcję portfeli efektywnych). Korzyści tych nie sposób kwestionować. Bernstein niejednokrotnie jednak fałszywie sugeruje, że wcześniej nie uświadamiano sobie znaczenia ryzyka i dywersyfikacji. Tymczasem praktycznie stosowaną dywersyfikację obserwujemy często w działaniach nawet prostych, niewykształconych ludzi. Zapewne najbardziej znanym przypadkiem jest wybór strategii połowów przez rybaków z Jamajki, wybierających kombinację słabych i pewnych oraz ryzykownych, lecz obfitych połowów⁴. Przypadek ten analizowany jest jednak w kategoriach teorii gier, jako gra przeciw naturze, a nie w kategoriach teorii finansów. Przewagą teorii Markowitza nad rozwiązaniem w ramach teorii gier jest ujęcie zagadnienia dywersyfikacji w eleganckie, ścisłe, matematyczne formuły, pozwalające na wprowadzenie pojęcia portfela efektywnego i jego zestawienie z funkcją użyteczności.

Z książki Bernsteina dowiadujemy się też, że Markowitz nie był pierwszy. W 1940 r., dwanaście lat przed Markowitzem, zagadnienie dywersyfikacji portfela przedstawił w zasadzie w ten sam sposób włoski matematyk Bruno de Finetti. Jego praca aż do 2006 r. nie została zauważona przez amerykańskich finansistów, zapewne dlatego, że dotyczyła ubezpieczeń, a nie inwestycji; przede wszystkim jednak została opublikowana po włosku. Jak widać, rynek wiedzy też nie jest rynkiem efektywnym.

Zalety teorii portfelowej Markowitza nie wystarczą jednak na jej obronę, zwłaszcza na obronę wyprowadzonego z niej modelu CAPM i powiązanych idei racjonalnych oczekiwań oraz rynku efektywnego (znamienne jest, że te pozytywne określenia nie oznaczają w istocie nic innego, jak działanie ludzi i funkcjonowanie rynków zgodne z „wytocznymi” klasycznej teorii finansów). Zauważmy, że według tych koncepcji możliwe jest prawidłowe wyznaczenie oczekiwanych stóp zwrotu z poszczególnych aktywów i ich wariancji. O ile takie założenie wydaje się być rozsądne wobec natury, gdzie występują pewne obiektywne, statystycznie powtarzalne zjawiska i prawidłowości, o tyle jest ono pozbawione podstaw w przypadku rynków finansowych. Wyrazem przyjętej postawy jest rozpatrywanie działalności na rynkach finansowych tak, jak gry przeciw naturze (jak w przypadku jamajskich rybaków kierujących

się wiecznymi prawami przyrody), a nie jak gry strategicznej. Bernstein pisze: „W dawnych czasach, kiedy większość działalności gospodarczej stanowiło łowiectwo, rybactwo i rolnictwo, pogoda była jedynym źródłem niepewności gospodarczej. Na pogodę nic się nie poradzi. W konsekwencji ludzie odwoływali się do modlitwy i zaklęć jako jedynej dostępnej formy zarządzania ryzykiem” (s. XIV). Wiemy już z przykładu jamajskich rybaków, że to nieprawda, a Bernstein zbyt hojnie przypisuje zasługi Markowitziowi. Następnie Bernstein trafnie rozpoznaje społeczny wymiar decyzji ekonomicznych, zwłaszcza na rynkach finansowych, oraz wskazuje, że wymaga to analizy w kategoriach teorii gier strategicznych wywodzącej się od von Neumanna. Są to trafne spostrzeżenia. Niestety jednak tego właśnie postulatu nie spełnia klasyczna teoria finansów, gdzie – jak zauważyliśmy – oczekiwania odzwierciedlają pewne, w zamyśle obiektywne zjawiska, a nie społeczne interakcje. Elementy gier strategicznych w teorii finansów (np. racjonalne – sic! – bąble lub kaskady informacyjne) pojawiają się dopiero wówczas, gdy porzucimy ideę rynku efektywnego, a uwzględnimy zjawiska będące przedmiotem behawioralnej teorii finansów.

Bernstein zdaje się nie dostrzegać, że klasyczna teoria finansów analizuje zagadnienie ryzyka statystycznego, a nie strategicznej niepewności wynikającej ze społecznego charakteru procesów zachodzących na rynkach finansowych; nie zmienia tego wrażenia nawet przytoczony cytat z rozmowy z W. Sharpem „Niebezpiecznie jest myśleć (...) o ryzyku jako liczbie” (s. 96). Nie dostrzegając tego „grzechu pierworodnego” teorii klasycznej, skłonny jest życzliwie traktować teorię behawioralną jako przydatny opis anomalii, które dzięki ich rozpoznaniu zostaną wyeliminowane i staną się żerem dla inwestorów-rekinów z wcześniejszego cytatu.

Bernstein w ten sposób nie tylko nie docenia destrukcyjnej wobec klasycznej teorii finansów siły teorii behawioralnej, lecz także niesłusznie chyba oczekuje od tej ostatniej konstruktywnych wniosków, które miałyby stanowić poprawkę do teorii klasycznej. Niestety, chociaż teoria behawioralna wskazuje systematyczne błędy w ludzkim rozumowaniu, to formułowane w jej ramach spostrzeżenia raczej nie tworzą spójnego modelu systematycznych błędów w wycenie aktywów. Teoria behawioralna nie jest aneksem do teorii klasycznej, nie jest także alternatywną teorią, która wyjaśniałaby kształtowanie się cen na rynkach finansowych. Ciekawe, że Bernstein zamieszcza w swojej książce cytaty, które wyraźnie prezentują takie stanowisko, na przykład wypowiedź Daniela Kahnemana: „Myślę, że modele behawioralne mogą być bardzo istotne w projektowaniu instytucji, ale nie jest jasne, czy ostatecznie mają one istotną moc objaśniającą ceny aktywów” (s. 31). Andrew Lo mówi o teorii behawioralnej, że jej odkrycia stanowią „zbiór anomalii, a nie teorię. Aby pokonać teorię, potrzeba innej teorii” (s. 61). Pomimo przytoczenia tych

⁴ Patrz: Straffin (2001, rozdz. 4).

sformułowań, Bernstein łudzi się, że uzupełniając teorię klasyczną obserwacjami z zakresu teorii behawioralnej można osiągnąć... nirwanę. Uważam zatem, że Bernstein stawia istotne pytania dotyczące teorii finansów, pobudza do myślenia, ale udziela niezadowolających lub niewłaściwych odpowiedzi.

Duża część książki nie stanowi w pełni czytelnego i bezpośredniego odniesienia do ogólnych kwestii omawianych powyżej, lecz jest prezentacją osób, instytucji i strategii inwestycyjnych. Większość rozdziałów jest owocem rozmów autora ze znakomitościami teorii finansów i amerykańskiego rynku finansowego.

Pierwsza, po wstępie, część książki pod tytułem „The Behavioral Attack” składa się z dwóch rozdziałów. Pierwszy z nich zawiera bardzo krótkie przedstawienie idei sformułowanych przez twórców teorii behawioralnej: Daniela Kahnemana, Amosa Tverskiego, oraz przynoszącej ponadprzeciętne wyniki działalności inwestycyjnej firmy, w którą zaangażowany jest sam Kahneman oraz Richard Thaler, inny prominentny przedstawiciel szkoły behawioralnej. Drugi rozdział stanowi raczej atak na teorię behawioralną niż przedstawienie istoty „ataku behawioralnego”. Bernstein utrzymuje, że na rozwiniętych rynkach inwestorzy wykorzystują wszystkie możliwości arbitrażu, nie pozostawiając możliwości systematycznego uzyskiwania ponadprzeciętnych zysków, które miałyby być skutkiem behawioralnych anomalii. W swoim rozumowaniu Bernstein zaciera jednak różnicę pomiędzy arbitrażem a spekulacją, której nieodłącznie towarzyszą ryzyko oraz niepewność i poprzez którą mogą się wyrażać zjawiska właściwe teorii behawioralnej.

Druga część książki („The Theoreticians”) zawiera aż siedem rozdziałów, stanowiących owoc rozmów autora z Paulem A. Samuelsonem, Robertem C. Mertonem, Andrew Lo, Robertem Schillerem, Williamem Sharpem, Harrym Markowitzem i Myronem Scholesem. Wśród tych twórców bądź kontynuatorów teorii klasycznej pewne zdziwienie może budzić osoba Schillera, który jako pierwszy już na początku lat 80. ubiegłego wieku zwrócił uwagę na nadmierną zmienność cen aktywów, a w 2000 r. opublikował książkę *Irrational Exuberance*. Jego spostrzeżenia stanowiły jedną z przesłanek teorii behawioralnej. Wydaje się jednak, że takie zestawienie odpowiada zamysłowi Bernsteina, który stara się pokazać ewolucję teorii klasycznej pod wpływem nowych idei. Z kolei w przypadku Scholesa włączenie go do grupy teoretyków (a nie kolejnej kategorii praktyków) można uzasadniać jedynie jego uhonorowanym nagrodą

Nobla udziałem w stworzeniu modelu Blacka-Scholesa-Mertona. Odpowiedni rozdział dotyczy wyłącznie aktualnej działalności biznesowej Scholesa. Znamienne, że ledwie mimochodem wspomniane jest w tym rozdziale bankructwo w 1998 r. funduszu LTCM, które zachwiało amerykańskim systemem finansowym. Scholes (i Merton) był właśnie partnerem LTCM (przypadek ten jest nieco szerzej omówiony w innym rozdziale).

Kolejna część książki („The Practitioners”) zawiera sześć rozdziałów. Nie we wszystkich daje się wyszczególnić centralną postać, wokół której osnuty byłby rozdział. Właściwie we wszystkich przypadkach jest to raczej przedstawienie pewnych praktycznych strategii inwestycyjnych. W wielu z nich rzeczywiście widać „ewolucję” teorii klasycznej, np. w rozdziale „CAPM II: The Great Ralph Dream Machine: We Don't See Expected Returns”. Bernstein z jednej strony przywołuje tu modele wieloczynnikowe, a z drugiej wika się w mgliste rozważania na temat strategii łączących zarządzanie strategiczne i taktyczne (a więc wychodzące z przesłanek rynku efektywnego i – przeciwnie – zaprzeczające im). Zdanie Bernsteina, że inwestorzy „poszukują wyższych zwrotów bez podejmowania nadmiernego ryzyka” (s. 177) w kategoriach CAPM oznaczałoby tyle samo, co zjeść ciasteczko i mieć ciasteczko. Osobiście za najciekawszy w tej części książki uznałem rozdział poświęcony wspaniałym osiągnięciom Davida Swensena w zarządzaniu funduszem Uniwersytetu Yale. Przypadek ten pokazuje wielkie możliwości, jakie stwarza inwestowanie na mniejszych rynkach, często towarowych, a nie *stricte* finansowych. Biorąc jednak pod uwagę zaangażowanie funduszu Uniwersytetu Yale w rynek nieruchomości, ciekawy byłby „suplement” do tego rozdziału, pokazujący wpływ obecnych zawirowań na rynku nieruchomości i rynku finansowym na wyniki funduszu.

Czwarta część książki (zawierająca tylko jeden, bardzo krótki rozdział) stanowi dość konwencjonalne zakończenie, po którym nie należy oczekiwać odpowiedzi na liczne pytania nasuwające się po lekturze.

W mojej ocenie książka Bernsteina często budzi zastrzeżenia, a tekst nie jest wolny od wad. Jej wielką zaletą jest jednak to, że pozwala zajrzeć za kurtynę najwyżej rozwiniętego w świecie amerykańskiego systemu finansowego i poznać idee współcześnie zajmujące teoretyków i praktyków. Kurtynę uchyla bywalec, znający najważniejszych aktorów. Warto skorzystać z tego zaproszenia.

Bibliografia

- Cieślak A. (2003), *Behawioralna ekonomia finansowa, modyfikacja paradygmatów funkcjonujących w nowoczesnej teorii finansów*, „Materiały i Studia”, nr 165, NBP, Warszawa.
- Straffin P.D. (2001), *Teoria gier*, Wydawnictwo Naukowe Scholar, Warszawa.

Andrzej Rzońca *Czy Keynes się pomylił? Skutki redukcji deficytu w Europie Środkowej*

Review of the book by Andrzej Rzońca, *Was Keynes Wrong? Consequences of Deficit Reduction in the Central Europe*

Wydawnictwo Naukowe „Scholar” oraz Fundacja Kronenberga Citi Handlowy, Warszawa 2007

*Tomasz Tokarski**

Dyskusje na temat restrykcyjności lub ekspansywności prowadzonej przez państwo polityki fiskalnej stanowią jeden z najważniejszych sporów zarówno wśród makroekonomistów z ośrodków akademickich, jak i wśród polityków gospodarczych na całym świecie. Zagadnienie to jest ważne głównie dlatego, że (w krótszej lub dłuższej perspektywie czasowej) decyduje o losach milionów gospodarstw domowych, dynamice wzrostu gospodarczego oraz kształtowaniu się zatrudnienia i bezrobocia w realnie funkcjonujących gospodarkach. Używając pewnego uproszczenia, można stwierdzić, że wśród polityków gospodarczych (bo niekoniecznie wśród ekonomistów z ośrodków akademickich) strony tego sporu dzielą się na zwolenników stymulowania zagregowanego popytu i wzrostu gospodarczego poprzez wydatki rządowe i podatki (nazywanych dalej umownie keynesistami) oraz zwolenników stosowania restrykcyjnej polityki fiskalnej (zwanych dalej niekeynesistami)¹. Keynesiści twierdzą,

że (szczególnie w okresach obniżenia aktywności gospodarczej) stymulowanie wielkości zagregowanego popytu może korzystnie oddziaływać (m.in. poprzez działanie efektów mnożnikowych) na wielkość produkcji i zatrudnienia, natomiast niekeynesiści uważają, iż dla utrzymania gospodarki na stabilnej ścieżce wzrostu gospodarczego najistotniejsze jest zachowanie dyscypliny fiskalnej. Problemom tym poświęcona jest też w głównej mierze praca Andrzeja Rzońcy.

Książka A. Rzońcy składa się z dwóch części. W pierwszej, teoretycznej części pracy (z tytułowanej „Konkurencyjne teorie”), scharakteryzowane są najistotniejsze modele makroekonomiczne opisujące oddziaływanie polityki fiskalnej na gospodarkę zarówno w krótkim, jak i (w pewnych przypadkach) w długim okresie. W drugiej części, pt. „Empiryczna weryfikacja teorii”, Autor przedstawia zarówno wyniki znanych z literatury, wybranych analiz empirycznych dotyczących skutków zacieśniania polityki fiskalnej, jak i własne analizy odnoszące się głównie do Węgier, Litwy, Estonii oraz Łotwy w okresie transformacji.

W teoretycznej części książki A. Rzońcy scharakteryzowane są następujące modele makroekonomiczne:

¹ Podział ten nie jest już tak jednoznaczny na gruncie makroekonomii akademickiej, gdyż np. uważany za keynesistę wybitny amerykański teoretyk wzrostu gospodarczego E. Domar analizował wpływ inwestycji (które mogą być również finansowane z wydatków rządowych) zarówno na popytową stronę gospodarki, jak i na jej możliwości podażowe. Podobnie R. Solow (twórca neoklasycznego modelu wzrostu gospodarczego), N. Mankiw, D. Romer i D. Weil uważani są za umiarkowanych keynesistów, mimo że w swoich opracowaniach dotyczących wzrostu gospodarczego (w przeciwieństwie do Keynesa) zajmowali się głównie podażową stroną gospodarki.

- Samuelsona,
- Hicksa-Hansena (model IS-LM),
- Mankiwa-Summersa,
- Silberta,
- Friedmana,
- Mundella-Fleminga,
- Ramseya-Cassa-Koopmansa,
- Blancharda,
- Lane’a-Perrotiego.

Wybór oraz charakterystyka modeli makroekonomicznych wykorzystanych w teoretycznej części pracy wydają się trafne i poprawne z punktu widzenia problematyki poruszanej w książce A. Rzońcy. Podziw musi też budzić różnorodność i liczba pozycji literatury, z których Autor korzystał w czasie pisania tej części pracy. Niewątpliwą słabością, w mojej ocenie, jest tu brak równań wykorzystywanych w przytaczanych przez Autora modelach równowagi krótko- i długookresowej. Często to, co słownie opisywane jest w opracowaniach ekonomicznych na kilku stronach, może być zapisane w postaci kilku mniej lub bardziej skomplikowanych równań, których rozwiązanie pozwala na wyciągnięcie wniosków natury ekonomicznej. Co prawda A. Rzońca (na s. 17) przytacza opinię P. Krugmana, że równania są „rusztowaniem ułatwiającym budowę intelektualnego gmachu; z chwilą gdy gmach ten zostanie wzniesiony do pewnego punktu, rusztowanie można usunąć, pozostawiając jedynie potoczny język”. Ja jednak nie zgadzam się z tą opinią, gdyż nie wszystkie zależności matematyczne da się opisać potocznym językiem (spróbujmy opisać w potocznym języku np. diagramy fazowe, które są często wykorzystywane w modelowaniu procesów wzrostu gospodarczego). Co więcej, sądzę, że we współczesnej makroekonomii zależności matematyczne (gdyż nie muszą to być tylko równania) są nie rusztowaniem, lecz fundamentem, na którym buduje się „gmach intelektualny”. Wyobraźmy sobie, co stanie się z budynkiem, jeśli po jego postawieniu usuniemy fundamenty. Gwoli uczciwości należy jednak stwierdzić, że w wielu przypadkach Autor odwołuje się do swoich wcześniejszych opracowań, w których można znaleźć również matematyczną stronę prezentowanych w książce rozważań teoretycznych. Sądzę jednak, że matematyczna strona omawianych zależności makroekonomicznych powinna się znaleźć (przynajmniej w formie aneksów) w książce A. Rzońcy, gdyż praca ta byłaby wówczas znacznie ciekawsza np. dla takich Czytelników, jak piszący te słowa.

W drugiej części (którą moim zdaniem nie do końca można nazwać empiryczną weryfikacją modeli teoretycznych) znajduje się zarówno przegląd wcześniejszych wyników analiz empirycznych dotyczących skutków zacieśnienia polityki fiskalnej, jak i rozważania A. Rzońcy na temat skutków takiej polityki fiskalnej na Węgrzech, na Litwie, w Estonii oraz na Łotwie w ciągu ostatnich dwóch dekad.

Przegląd wcześniejszych wyników badań empirycznych uważam za ciekawy i pouczający dla Czytelników książki. Słabością tego przeglądu jest jednak to, że A. Rzońca rzadko wspomina o tym, jakimi metodami statystycznymi posługiwali się cytowani przez niego autorzy. Stąd też Czytelnicy mają stosunkowo mało podstaw do wyrobienia sobie zdania na temat wiarygodności przytaczanych w tej części pracy wyników analiz empirycznych.

Przechodząc do empirycznych analiz książki A. Rzońcy (rozdziały XIV–XVII), należy przede wszystkim zwrócić uwagę na bogactwo danych statystycznych, które wykorzystano w tej części omawianej pracy. Uzasadniając tezę o niekeynesistowskich skutkach zacieśnienia polityki fiskalnej w czterech wspomnianych uprzednio krajach Europy Środkowo-Wschodniej, Autor analizuje wpływ ograniczenia wydatków budżetowych państwa lub wzrostu jego dochodów głównie na prywatną konsumpcję, eksport, inwestycje oraz tempo wzrostu gospodarczego (mierzone zwykle tempem wzrostu PKB). Analizy te zazwyczaj prowadzone są w roku poprzedzającym ograniczenie deficytu budżetowego, w okresie zacieśnienia fiskalnego oraz w rok po procesie rozważanym w książce A. Rzońcy. Prowadzone przez Autora rozważania na pozór wydają się przekonujące. Niemniej jednak mogą być obciążone kilkoma istotnymi mankamentami. Do mankamentów tych zaliczam głównie to, że:

- Autor w zasadzie abstrahuje od koordynacji polityki fiskalnej i monetarnej (bodaj w jednym tylko miejscu w liczącej ponad 300 stron książce wspomina o tym, że instrumenty polityki monetarnej nie mogą amortyzować skutków polityki fiskalnej, co wydaje się tezą dość kontrowersyjną).
- Wydaje się także, że na wielkość realizowanego w gospodarce produktu wpływa zarówno to, co dzieje się po popytowej stronie gospodarki, jak i to, co ma miejsce po jej stronie podażowej (chyba jedyne występujące w książce A. Rzońcy analizy dotyczące podażowej strony gospodarki odnoszą się do wpływu płac i kursu walutowego na konkurencyjność eksportu). Analizując skutki zacieśniania polityki fiskalnej Autor odnosi się głównie do zmian po stronie produkcji, inwestycji, prywatnej konsumpcji oraz eksportu. Analizy takie mają głównie charakter analiz popytowych (czyli w jakiejś mierze również keynesistowskich). Sądzę, że rozważając skutki zacieśnienia polityki fiskalnej, należałoby sprawdzić, jak to zacieśnienie oddziałuje nie tylko na tempo rzeczywistego PKB, ale również na różnicę między tempem wzrostu rzeczywistego PKB a potencjalnego PKB². W ten sposób oddzieliłoby się skutki wpływu zacieśnienia polityki fiskalnej na popytową stronę gospodarki od tempa wzrostu potencjalnego PKB, które szczególnie w pierwszych latach transformacji systemowej w dużym

² Inną sprawą jest to, czy dla gospodarek transformacji (szczególnie na początku procesu transformacji systemowej) możliwe jest wiarygodne statystycznie oszacowanie ścieżki produktu potencjalnego.

stopniu zależało zarówno od czynników zewnętrznych, jak i od zmian rzeczowej struktury popytu, struktury produktu oraz zmian na rynku pracy. Analizując zaś oddziaływanie zacieśnienia polityki fiskalnej jedynie na rzeczywisty PKB, nie sposób oddzielić efekty tej polityki od uwarunkowań zewnętrznych, jak też uwarunkowań wewnętrznych niezwiązanych bezpośrednio z polityką fiskalną.

- Sądzę również, że od zestawiania ze sobą pewnych par danych statystycznych (zazwyczaj w przedziałach 3–4-letnich) znacznie bardziej użytecznym narzędziem analitycznym są modele ekonometryczne złożone z kilku (lub może kilkunastu) równań dla każdej z gospodarek analizowanych w pracy A. Rzońcy. Modele takie można byłoby skonstruować, np. wykorzystując kwartalne dane statystyczne dla tych gospodarek, poczynwszy – powiedzmy – od 1994 lub 1995 r. Konstrukcja i analiza parametrów takich modeli ekonometrycznych miałyby przynajmniej kilka zalet w stosunku do analiz jakościowych prowadzonych w pracy A. Rzońcy. Po pierwsze, pozwalałaby na precyzyjniejszy ilościowo opis sprzężeń zwrotnych pomiędzy podstawowymi zmiennymi makroekonom-

icznymi (nie tylko związanymi z instrumentami polityki fiskalnej) na Węgrzech, na Litwie, w Estonii i na Łotwie. Po drugie, modele ekonometryczne pozwalają na oszacowanie rozkładów opóźnień pomiędzy zmiennymi (na gruncie prowadzonych przez A. Rzońcę analiz jakościowych długość tych opóźnień może być dobierana dość arbitralnie, Czytelnik musi tu wierzyć Autorowi na słowo). Po trzecie, wykorzystanie takich modeli umożliwiłoby również szacunkowe oddzielenie skutków zacieśnienia polityki fiskalnej od skutków zdarzeń niezależnych od owego zacieśnienia. Po czwarte, modele ekonometryczne pozwalają również na symulacje zachowań najważniejszych zmiennych makroekonomicznych w przypadku, w którym zaniechano by restrykcyjnej polityki fiskalnej (symulacje takie są, moim zdaniem, raczej niemożliwe na gruncie analiz jakościowych).

Podsumowując, sądzą, że książka A. Rzońcy *Czy Keynes się pomylił?* jest ważnym głosem w prowadzonej w Polsce dyskusji na temat zacieśniania bądź rozluźniania polityki fiskalnej. Niemniej jednak nie przekonuje mnie, że odpowiedź na postawione w tytule pytanie powinna być twierdząca.