

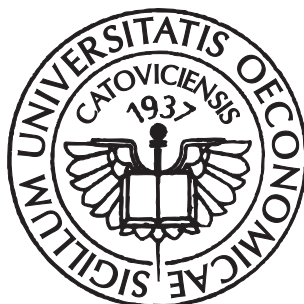
**STATYSTYCZNO-DYNAMICZNE
MODELE ZARZĄDZANIA
RYZYKIEM EKONOMICZNYM
PROGNOSTYCZNE UWARUNKOWANIA
W GOSPODARCE I ANALIZACH SPOŁECZNYCH**

Studia Ekonomiczne

**ZESZYTY NAUKOWE
WYDZIAŁOWE
UNIWERSYTETU EKONOMICZNEGO
W KATOWICACH**

**STATYSTYCZNO-DYNAMICZNE
MODELE ZARZĄDZANIA
RYZYKIEM EKONOMICZNYM
PROGNOSTYCZNE UWARUNKOWANIA
W GOSPODARCE I ANALIZACH SPOŁECZNYCH**

**Redaktor naukowy
Włodzimierz Szkutnik**



Katowice 2012

Komitet Redakcyjny

Krystyna Lisiecka (przewodnicząca), Anna Lebda-Wyborna (sekretarz),
Halina Henzel, Anna Kostur, Maria Michałowska, Grażyna Musiał, Irena Pyka,
Stanisław Stanek, Stanisław Swadźba, Janusz Wywiał, Teresa Żabińska

Komitet Redakcyjny Wydziału Ekonomii

Stanisław Swadźba (redaktor naczelny), Magdalena Tusińska (sekretarz),
Teresa Kraśnicka, Maria Michałowska, Celina Olszak

Rada Programowa

Lorenzo Fattorini, Mario Glowik, Miloš Král, Bronisław Micherda,
Zdeněk Mikoláš, Marian Noga, Gwo-Hsiung Tzenga

Recenzent

Paweł Dittmann

Redaktor

Magdalena Bulanda

Skład tekstu

Urszula Grendys

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2012

ISBN 978-83-7246-770-6

ISSN 2083-8611

Wszelkie prawa zastrzeżone. Każda reprodukcja lub adaptacja całości bądź części
niniejszej publikacji, niezależnie od zastosowanej techniki reprodukcji,
wymaga pisemnej zgody Wydawcy

WYDAWNICTWO UNIWERSYTETU EKONOMICZNEGO W KATOWICACH

ul. 1 Maja 50, 40-287 Katowice, tel. 32 25 77 635, fax 32 25 77 643
www.ue.katowice.pl, e-mail: wydawnictwo@ue.katowice.pl

SPIS TREŚCI

WSTĘP	7
Włodzimierz Szkutnik: MODELOWY KONTEKST KONTROWERSYJNEJ AKSJOLOGII RACJONALNYCH OCZEKIWAŃ W EKONOMII	11
Summary	24
Alicja Wolny-Dominiak: UOGÓLNIONA METODA TARYFIKACYJNA KLASY MBP W UBEZPIECZENIACH KOMUNIKACYJNYCH	25
Summary	36
Monika Dyduch: WYCENA PRODUKTU STRUKTURYZOWANEGO NA PRZYKŁADZIE LOKATY STRUKTURYZOWANEJ STABILNA ŻŁOTÓWKA	39
Summary	54
Aleksandra Szpulak: PROGNOZOWANIE W ZARZĄDZANIU AKTYWAMI OBROTOWYMI PRZEDSIĘBIORSTWA	55
Summary	70
Elżbieta Sojka: ANALIZA PRZESTRZENNO-CZASOWA RYNKU PRACY W POLSCE ZE SZCZEGÓLNYM UWZGLĘDNIENIEM PROBLEMU BEZROBOCIA	71
Summary	87
Jan Acedański: RYNEK AKCJI A KOSZTY WAHAŃ KONIUNKTURALNYCH W POLSCE	89
Summary	101
Maria Balcerowicz-Szkutnik: PROBLEMY STARZENIA RYNKU PRACY – ANALIZA STATYSTYCZNO-DEMOGRAFICZNA DLA WYBRANYCH PAŃSTW UE	103
Summary	115

Maciej Pichura: WPŁYW ASYMETRII ZMIENNOŚCI CEN NA EFEKTYWNOŚĆ INWESTYCJI NA RYNKU KAPITAŁOWYM	117
Summary	130
Wanda Ronka-Chmielowiec: UBEZPIECZENIA KOMUNIKACYJNE JAKO PODSTAWOWE PRODUKTY UBEZPIECZENIOWE SPRZEDAWANE NA POLSKIM RYNKU UBEZPIECZENIOWYM	131
Summary	146
Włodzimierz Szkutnik: WPŁYW RYZYKA STOPY BAZOWEJ I MORALNEGO HAZARDU NA CENĘ I WYPŁATĘ OBLIGACJI CAT	147
Summary	175
Włodzimierz Szkutnik: REGRESYJNE I DYNAMICZNO- -STATYSTYCZNE MODELOWANIE ZMIENNOŚCI W PRAKTYCE OCENY RYZYKA INWESTYCJI KAPITAŁOWYCH I UBEZPIECZENIOWYCH	177
Summary	191
Bartosz Lawędziak: OCENA SEKURYTYZACJI JAKO PROCESU ODDZIAŁYWUJĄCEGO NA STABILNOŚĆ FINANSOWĄ SYSTEMU BANKOWEGO	193
Summary	210

WSTĘP

W niniejszym zeszycie Studiów Ekonomicznych pt. „Statystyczno-dynamiczne modele zarządzania ryzykiem ekonomicznym. Prognostyczne uwarunkowania w gospodarce i analizach społecznych” prezentowane są tematy aktualnie odnoszące się do zjawisk obserwowanych na rynkach kapitałowym, ubezpieczeniowym, pracy i w sferze społecznych oddziaływań zmiennych uwarunkowań ekonomicznych.

Rozpoczynając od ogólnie ujętego modelowego kontekstu kontrowersyjnej aksjologii racjonalnych oczekiwań w ekonomii, poszczególne tematy rozpatrzone przez autorów dotyczą konkretnych, ścisłych zagadnień, w których w różnych aspektach przejawia się ryzyko ekonomiczne i zarządzanie tym ryzykiem z prognostycznym postrzeganiem uwarunkowań gospodarczych i społecznych.

Podstawą opracowań poszczególnych artykułów jest koncepcja historycznego postrzegania zjawisk oraz postawa właściwa dla przyjmowanych przez badaczy od ponad 50 lat koncepcji umocowanych w teorii racjonalnych oczekiwań. Idea tej koncepcji przeniknięta jest deterministyczną idealizacją rzeczywistości i jest zgodna z tymi nurtami w teorii ekonomii, które odnoszą się do adekwatnych względem niej aksjologicznych postaw i zachowań społecznych, a w szczególności ekonomicznych. Istotna w tym ujęciu jest dostrzegalna dychotomia między brakiem jednolitej teorii w ekonomii jako nauki społecznej a różnorodnością koncepcji i odniesień do innych dziedzin nauk, w tym humanistycznych i technicznych. Różnorodność postrzegania określonych zjawisk społecznych, w tym m.in. demograficznych oraz przeplatanie się określonych postaw o charakterze filozoficznym i politycznym jest źródłem nieporozumień oraz twórczych nowych trendów w rozwoju ekonomii jako nauki.

Tematy poszczególnych artykułów mają charakter formalny i analityczny. Autorzy podjęli tematy odnoszące się do współczesnych problemów społecznych oraz gospodarczych z uwzględnieniem przenikającego je ryzyka kapitałowego, ubezpieczeniowego demograficznego i ogólnie określając, finansowego.

W grupie opracowań tematów ubezpieczeniowych zwraca uwagę uogólniony schemat taryfikowania w ubezpieczeniach komunikacyjnych, które są analizowane także w perspektywie produktów sprzedawanych na polskim rynku kapitałowym. Do tej klasy tematów należy zaliczyć także sekurytyzację niemogącą, jak dotąd, zająć znaczącej pozycji w literaturze z zakresu ubezpieczeń. Sekurytyzacja, pojmowana również jako forma finansowania przedsięwzięć, oceniana jest w niniejszych Studiach jako proces oddziałujący na stabilność finansową systemu bankowego. Natomiast z innej strony jest postrzegana w procesie zarządzania sekurytyzacją katastrofalnego ryzyka ubezpieczeniowego. W tym drugim sformułowaniu została dokonana symulacja cen obligacji katastroficznych, których celem jest finansowanie skutków katastroficznych wydarzeń.

Wycena w systemie bankowym produktów strukturyzowanych, zyskujących na znaczeniu w procedurach wynikających ze struktury zarządzania w finansach, została zaprezentowana z zastosowaniem oryginalnej metody transformacji falkowej. Ryzyko kapitałowe zostało natomiast ocenione w opracowaniu asymetrycznych miar zmienności cen rozpatrzonych w kontekście ich wpływu na efektywność inwestycji na rynku kapitałowym. Zmienność jako taka, będąca m.in. w wycenach produktów rynku kapitałowego źródłem istotnych problemów, np. formalnych, została w ogólnym ujęciu uwzględniona oraz przedstawiona w sformalizowanym modelu regresyjnym i dynamiczno-statystycznym. Zastosowanie tych modeli wynika z potrzeb praktyki oceny ryzyka inwestycji kapitałowych i ubezpieczeniowych. Problemy podjęte w przedstawionej powyżej grupie tematycznej są także ogólnie potraktowane w artykule dotyczącym reakcji rynku akcji na koszty wahań koniunkturalnych. Ten ważny temat łączy ekonomicznie spektakularny charakter gospodarki rynkowej z jego elementem prezentowanym przez rynek akcji. Przedstawione tu podejście wywodzi się z modeli dynamiczno-stochastycznych gospodarki otwartej, znanych w kręgach bankowych analityków ze względu na ich swoistą złożoność formalną i teorii koncepcyjnie wymagającej solidnej znajomości podstaw ekonomicznych, a także ze względu na ich kontrowersyjność metodologiczną.

Pragmatyczne podejście do prognozowania z aplikacją w zarządzaniu zostało przedstawione w kontekście aktywów obrotowych przedsiębiorstwa. Stanowi to interesujące ze względów praktycznych pole analiz bezpośrednio przydatnych dla podejmowania decyzji w przedsiębiorstwie, zawężone do jego aktywów obrotowych.

Analizy społeczne z uwzględnieniem rynku pracy zostały zaprezentowane w układzie przestrzenno-czasowym na przykładzie Polski w aspekcie bezrobocia. Natomiast pierwiastek demograficzny odniesień analiz na rynku pracy ukazany w kontekście starzenia rynku pracy, został przedstawiony dla wybranych krajów Unii Europejskiej.

Wszechstronnie zobrazowane badania przedstawione w artykułach są tematem badań własnych autorów, finansowanych ze środków badań statutowych. Ich autorzy reprezentują Katedrę Metod Statystyczno-Matematycznych w Ekonomii Uniwersytetu Ekonomicznego w Katowicach oraz Katedrę Ubezpieczeń i Katedrę Analiz i Prognoz Gospodarczych Uniwersytetu Ekonomicznego we Wrocławiu. Część badań przeprowadzili młodzi pracownicy nauki przygotowujący rozprawy doktorskie i habilitacyjne.

Włodzimierz Szkutnik

Włodzimierz Szkutnik

MODELOWY KONTEKST KONTROWERSYJNEJ AKSJOLOGII RACJONALNYCH OCZEKIWAŃ W EKONOMII

1. Model formalny w tle przeszłych i współczesnych teorii ekonomicznych

W analizach ekonomicznych dokonuje się z reguły próby wyjaśniania aktualnego i mającego nastąpić w czasie przyszłym stanu realnej ekonomii regionu, kraju oraz globalnego układu. W zależności od ujęcia tematu badań można uzyskać obraz rzeczywistości względnie dokładnie do niej dopasowany. Złożoność realnej ekonomii wymaga zawsze posługiwania się wyrafinowanym aparatem opisowo-formalnym wypracowanym przez pojawiające się sekwencyjnie na przestrzeni dziejów teorie ekonomiczne. Każda z nich została stworzona „na potrzeby” danego etapu rozwoju gospodarczego i społecznego. Tak było z klasyczną teorią liberalną A. Smitha, J.S. Milla, D. Ricardo, teorią K. Marksa materialistyczną w swej dialektyce z idealistycznymi odniesieniami do Hegla, neoklasyczną teorią ekonomii Marshalla, Walrasa i innych, teorią Keynesa i wcześniejszą od niej, lecz nieujawnioną w języku angielskim, teorią Kaleckiego, jej licznymi odmianami oraz współczesnymi koncepcjami, w tym heterodoksyjnej myśli ekonomicznej (instytucjonalizm i postkeynensizm). Wybiegały one w przyszłość, starając się rozwikłać trapiące ludzkość fundamentalne problemy. Niezmiennie istotnym obszarem podejmowanym w analizach pozostaje tempo i granice wzrostu gospodarczego, problem równowagi w rozwoju gospodarczym oraz cykle koniunkturalne trapiące teoretyków ekonomii. Istotnym zagadnieniem podejmowanym nie od dzisiaj jest włączenie do tych teorii elementów subiektywizmu w rozstrzyganiu, jakie decyzje i sposób działania powinny być zastosowane dla osiągnięcia zamierzonych celów. Zasadniczy problem tkwił jednak zawsze w rozstrzygnięciu, czy pragmatyczne działania w wielowymiarowej przestrzeni historycznej, geograficznej, kulturowej, instytucjonalnej i politycznej mogą być siłą motoryczną w rozwiązywaniu splecionych oraz uwikłanych kwestii ekonomicznych i społecznych w wymiarze krajowym i globalnym. W obecnym czasie zbieżność różnorodnych zjawisk czy też zbieg okoliczności wywołuje lawinowe reakcje w skali globalnej. Może to

się nasilać, co stwarza olbrzymie pole reakcji dla pragmatycznej koncepcji podejmowania prób sterowania rozwojem międzynarodowych powiązań gospodarczych i globalizacji, poprzez umocnienie i koordynację funkcjonowania instytucji oraz regulacji międzynarodowych. Wychodząc z założeń koniecznej identyfikacji, tzw. *global governance*, czyli intensyfikacji programów sterowania globalizacją, być może rozwiązanie współczesnych problemów jest możliwe. Należy tu zwrócić uwagę na problemy, w których *global governance* ma duże szanse zaistnienia, a więc tam, gdzie występuje duża zbieżność interesów poszczególnych krajów i ich grup. Oprócz wspomnianego tempa oraz granic wzrostu gospodarczego wyraźne odniesienie *global governance* istnieje w zakresie wypracowania wartości i ich kulturowych implikacji dla procesów rozwojowych, modernizacji instytucjonalnej globalizacji *versus* narastający chaos i brak koordynacji, integracji regionalnej i jej sprzężenia z globalizacją. Ponadto znaczące miejsce powinno przynależeć dla pozycji i roli organizacji pozarządowych. Niemniej ważne aspekty powinny być uwzględniane w kształtowaniu środowiska przyrodniczego i regulacji kwestii wyczerpujących się zasobów naturalnych. Wreszcie istotne miejsce w tworzeniu *global governance* powinny mieć sprawy społeczne dotyczące procesów demograficznych, biedy, nędzy i nierówności społecznej, gospodarki oraz społeczeństwa opartego na wiedzy, postępu naukowo-technicznego, ewolucji sieci i jej gospodarczych konsekwencji, konflikty i bezpieczeństwo.

W niniejszym artykule podstawą rozważań jest przyjmowana przez badaczy od ponad 50 lat koncepcja racjonalnych oczekiwań będąca ujmującym rozstrzygnięciem postaw i zachowań społecznych oraz ekonomicznych w rozwiązywaniu w uwzględnianych w modelach istotnych dla ekonomii czasów nowożytnych. Krytyka tej koncepcji według R. Frydmana i M. Goldberga, zaprezentowana w ostatnich dwudziestu latach i akceptowana przez współpracującego z M. Friedmanem w rozwiązywaniu kwestii teoretycznych monetaryzmu E. Phelpssem, laureatem nagrody Nobla z 2006 r., nie dostarcza przekonujących argumentów na rzecz odrzucenia teorii racjonalnych oczekiwań. Jest jednak istotnym wkładem w postrzeganie problemów nierozwiązanych w ekonomii, nazywanej przez tych badaczy ekonomią niepewności. Należy jeszcze raz odnieść się do stanowiska o niejednoznaczności rozwiązań panujących w ekonomii od początku czasów, gdy zaczęto objaśniać procesy gospodarcze. Ekonomia jako nauka społeczna nie ma swej jednolitej teorii, co jest z jednej strony pozytywne, a z drugiej wyzwala lawinę różnych koncepcji i odniesień do innych dziedzin, z których dla uzasadnienia swych spekulacji logicznych, często de-skryptywnych, czerpie wiele rozwiązań. Zauważalna dywersyfikacja ujęć problemów społecznych, w tym gospodarczych, a także przenikanie się różnych dyscyplin jest właściwe wszystkim nurtom ekonomii na przestrzeni dziejów. Istotny jest tu splot tych teorii z polityką, które łączy pomost filozoficznego egzystencjalnego postrzegania bytu człowieka.

Wraz z rozwojem gospodarki kapitalistycznej, co jest notowane w doktrynach ekonomicznych, następowało rozpraszanie się kapitału. Ta dekoncentracja wpływała i wpływa na zmianę struktury własności z kapitałowej na akcyjną, chociaż okresowo występuje swoisty regres i powrót do silnej koncentracji właścicielskiej i skupiania się kapitału w jednostkowym wymiarze. Rola kapitalistów została zastąpiona zarządzającymi kapitałem, którzy bywają opatrzenie postrzegani jako właściciele majątku, którym zarządzają, a często mentalnie sami się kreują do takiej roli, poprzez przyjmowaną postawę społeczną i sposób zachowania. Powyższe bardzo skrótowe wprowadzenie w tło teorii ekonomicznych nie uwzględnia podstawowych relacji i kategorii ekonomii. Podstawowe z nich w analizach ekonomicznych, relacje cen, popytu i podaży, konkurencji, koncentracji kapitału i centralizmu są istotne w kwestiach związanych z transformacją kapitału, co nieodłącznie spleta się z kategoriami ekonomii: z zyskiem, cenami i płacami równowagi, konkurencją, inflacją, bezrobociem czy też stanem permanentnej wektoryzacji gospodarki kapitalistycznej utrzymującej w pozycji wzrostowej zysk z prowadzonej działalności. Istotne w kontekście wyjaśniania wzrostu są teorie ludnościowe. Inny z kolei nurt rozważań, kategoria wartości, w pośredni sposób prowadzi do użyteczności, wiąże się też z dobrobytem oraz ogólnym poziomem cen i ich relacji. W zależności od wyznawanego poglądu czy postrzegania realnych zjawisk gospodarczych i społecznych akcent w teoriach ekonomicznych był położony na inną relatywną cechę czy kategorię ekonomii. Istotną kwestią była akumulacja kapitału. Klasyczna teoria wartości wydawała się nieadekwatna do opisu stosunków ekonomicznych (kosztów produkcji). Według neoklasyków (Walras) wysokie koszty poniesione dla wytworzenia danego dobra nie muszą konieczności prowadzić do wysokich cen. Wyprowadzona przez nich użyteczność krańcowa uzależnia pojęcie wartości od użyteczności (zapotrzebowanie na dobra) i odnosi się do przyszłości, a nie do przeszłości. Zatem popyt wpłynie na ustalenie ceny, co dotknie nieświadomych tego producentów. Szkoła użyteczności krańcowej przyznawała, że czynniki wytwórcze reprezentują pewne wartości, ale ich zakres był wyznaczany przez użyteczność krańcową konsumpcji dóbr krańcowych wytwarzanych przez te czynniki. Klasycy posługiwali się użytecznością całkowitą, co było przyczyną powstania pewnych paradoksów w ich rozumowaniu. Znamienne jest, że nie stosowali określenia użyteczności, lecz „intensywności satysfakcji”. Według nich użyteczność była zjawiskiem psychologicznym z niewyspecjalizowaną jednostką miary. Jako cecha dóbr krańcowych, czyli konsumpcyjnych, użyteczność powinna być mierzona (według Jevonsa: użyteczność nabyta). Zasada malejącej użyteczności krańcowej (wypracowana później) według neoklasyków oznaczała, że w miarę jak rośnie konsumpcja dobra, jego użyteczność krańcowa maleje. Było to oparte na założeniu, że niezależnie od tego, czym jest użyteczność krańcowa, można ją mierzyć. Neoklasycy zakładali tę mierzalność użyteczności w liczebnikach głównych.

Jevons i Walras używali funkcji użyteczności i zakładali w idealizacji tego pojęcia, że ilość konsumowanego dobra oraz ilość użyteczności są w sposób ciągły podzielne. Uznając nierealność tego założenia, uwzględniali niepodzielność użyteczności, co powodowało nieciągłość funkcji użyteczności. W gruncie rzeczy obaj razem z Mengerem rozwinęli techniczny aparat badawczy nowoczesnej ekonomii i wywarli znaczący wpływ na myśl ekonomiczną dwudziestego wieku. Jednak to Walras i Marshall pretendują do miana twórców neoklasycznego nurtu w ekonomii, chociaż sam termin ekonomia neoklasyczna stworzył T. Veblen w 1990 r., przewrotnie, dla wyrażenia dezaprobaty wobec podejścia A. Marshalla. Zarówno Walras, jak i Marshall różnią się zasadniczo w swych poglądach. Wielu teoretyków ekonomii, jak np. Colander, uważa, że w literaturze istnieją liczne sprzeczności dotyczące klasyfikacji ekonomii na klasyczną i neoklasyczną, a także liczne polemiki z jej myślą reprezentowaną głównie przez Mengera, Veblena i Walrasa. Właśnie do tej myśli ekonomii neoklasycznej nawiązują także twórcy teorii racjonalnych oczekiwań.

Nie wnikając w dalsze szczegóły postrzegania przez neoklasyków użyteczności, należy tu stwierdzić, że niektóre z nich (jeśli nie większość) nie wytrzymało krytyki czasu i obecnie w nowoczesnej mikroekonomii uznaje się stosunki komplementarności oraz substytucji dóbr jako funkcji całkowitej użyteczności. Natomiast Gossen, z którym nie zgadzał się m.in. Walras, głosił, że konsumenci maksymalizują swoją całkowitą użyteczność poprzez dokonywanie zakupów, aż ostatnia jednostka pieniężna wydatkowana na jakieś dobro osiągnie taką samą użyteczność krańcową, jak ostatnia jednostka wydatkowana na każde inne dobro.

Tylko Walras potrafił zarówno ustalić stosunek między użytecznością a popytem, jak i wykazać, że fundamentalnym czynnikiem, który się kryje za popytem, jest użyteczność krańcowa. Prowadzi to do neoklasycznej teorii wartości.

Nie odnosząc się bliżej i nie opowiadając się specjalnie za żadnym ze stanowisk prezentowanych przez przedstawicieli poszczególnych szkół względem użyteczności klasyków (szczególnie J. Milla i D. Ricardo), należy wyraźnie podkreślić, że trendy uprawiane przez pierwszą generację teoretyków marginalizmu przetrwały do chwili obecnej. Rozwijane w czasach współczesnych abstrakcyjne modele z imponującym eksponowaniem technik matematycznych i statystycznych są pokłosiem ujęcia problemów ekonomicznych przez przeszłe szkoły. Jest przy tym oczywiste, co przyznaje wielu wybitnych ekonomistów, że ujęcia wypracowane przez teoretyków marginalizmu drugiej generacji, dotyczące produkcji, kosztów, cen czynników wytwórczych i podziału dochodu, wynikają ze słabych stron mikroteorii w ujęciu Walrasa oraz jego modelu ogólnej równowagi.

2. Niepewność, ryzyko i racjonalność w wymiarze zmienności procesów gospodarczych

Zmienność zjawisk fizycznych, społecznych i gospodarczych jest emanacją niepewności oraz ryzyka jako pojęć, których znaczenie w ekonomii przekroczyło wszelkie oczekiwania. Pojęcia niepewności i ryzyka wprowadzone głównie przez przedstawicieli neoklasycznej szkoły ryzyka w ekonomii są wygodnym sposobem objaśniania wszystkiego, czego nie można uzasadnić w logicznie uporządkowanym dedukcyjnym systemie pojęć pierwotnych w znaczeniu matematycznym.

W modelowaniu procesów gospodarczych, w tym inwestycyjnych w zakresie ubezpieczeniowym oraz kapitałowym, niepewność i ryzyko są w szczególności sposoby regulowane i identyfikowane przez zmienność cen oraz wartości aktywów, w szczególności akcji i opcji. Ogólnie rzecz biorąc, zmienność nie zawsze ma konotacje sformalizowane aksjomatyką probabilistyczną i wtedy przyjmuje się, że warunki przejawiania się zjawiska ekonomicznego mają charakter niepewności. W przeciwnym przypadku występuje ryzyko, którego ocena jest podstawą decyzji zarządczych, których celem jest zniwelowanie negatywnych skutków mogących się zrealizować, gdyby ryzyko napotkało na sprzyjające warunki dla wytworzenia się procesu strat. Warunki te mogą być implikowane przez określony spłot okoliczności mikro- lub makroekonomicznych wyzwolonych m.in. przez politykę pieniężną, inwestycyjną, monetarną itp. Kontekst ubezpieczeniowy, co należy tu zaakcentować, wpleciony w decyzje podejmowane na rynku kapitałowym i szerzej finansowym może mieć wiele odmian, co wpływa na jego pośrednie oddziaływanie na wiele niezidentyfikowanych czynników mających wpływ na gospodarkę jako całość.

Zmienność rozpatrywana w dalszej części jest ujęta z jednej strony w kontekście wymogów modelowych o charakterze statystycznym i dynamicznym, a drugiej strony ma ona odwzorowywać niejednoznaczną oraz niekomplementarną naturę niepewności i ryzyka. Literatura na ten temat jest bardzo obszerna i nie jest celem prowadzonych tu rozważań zajmowanie się różnymi odmianami niepewności, a szczególnie ryzyka. Wynika to po części z nierozwikłanych, niekiedy niespójnych treści merytorycznych tych pojęć widzianych przez pryzmat prowadzonych w różnych ujęciach rozważań deskryptywnych i stricte formalnych.

Wyjaśnianie procesów gospodarczych, reakcji rynków, w tym kapitałowego, które transmitowane do gospodarki wpływają na nastroje i zachowania konsumentów, implikuje na zasadzie reakcji wstecznej procesy gospodarcze

i następnie wpływa na behawioryzm rynków, a właściwie jego podmiotów. Przebiega to w czasie, zatem dynamika oraz charakter tych procesów prowadzą do wielorakich ujęć procesów makroekonomicznych, które są oparte ogólnie rzecz biorąc, na teorii racjonalnych oczekiwań. Teoria ta wywarła nieoceniony wpływ na objaśnianie polityki gospodarczej w wymiarze światowym.

2.1. Racjonalność w przestrzeni gospodarki globalnej

Zagadnienie racjonalności we wszelkich formach działalności, szczególnie działalności gospodarczej człowieka, bywa podstawową ideą prowadzącą ludzkość ku rozwojowi. W celu oceny racjonalności procesów gospodarczych w znaczeniu globalnym, co jest nieco odmiennym pojęciem od tego pojmowanego w teorii racjonalnych oczekiwań, ale od którego warto wyjść w zawartych tu rozważaniach, należałoby skupić uwagę na samej definicji pojęcia dobra globalnego. W ocenie takiej należy uwzględnić punkt odniesienia, jakim jest analiza celów stawianych przez człowieka w swej działalności. Dlatego nieodzowne jest rozdzielenie dobra od niebezpieczeństwa w tej działalności. Nie zidentyfikowanie tych obszarów prowadzi zwykle do niewłaściwego klasyfikowania działań. Skupienie następuje na jednych procesach, które są traktowane priorytetowo, przy jednoczesnym niedocenianiu innych. W takim wypadku mamy do czynienia z oczywistą niepewnością efektów. W znaczeniu zmienności można mówić wtedy o jej zupełnej dowolności, a w sensie antycypacji efektów działań występuje całkowita niewiedza.

Znamienne jest, że w realnej ekonomicznej, społecznej i środowiskowej przestrzeni działań, występuje swoiste wymuszenie realizacji racjonalności. Racjonalność w tym aspekcie będzie miała jednak charakter wycinkowy na tle racjonalności globalnej; mianowicie będzie miała charakter irracjonalny, co może dowodzić (gdyż nie ma tu całkowitego przekonania, że tak będzie zawsze) o niekompletności procesów globalizacji. Racjonalność dla jednostki nie jest bowiem równoznaczna z racjonalnością działań egzystencjalnych, co przekonuje raczej do powyższego stwierdzenia.

Niekompletność procesów globalizacyjnych, na co zwraca uwagę wielu autorów, także polskich, prowadzi do bardzo wielu zaprzeczeń i wręcz nieracjonalnych ocen procesów obserwowanych w związku z globalizacją, co nie oznacza wcale, że sama globalizacja jest tego przyczyną. Jest to stwierdzenie znaczące samo w sobie i mające nośny intelektualnie charakter.

Powyższa hipoteza w badaniach mechanizmów społeczno-ekonomicznych sprowadza się do identyfikacji przyczyn i skutków tej niekompletności procesów globalizacji. Zwraca tu uwagę oczywiste wręcz rozdzielnie się sfery makro- i mikroekonomicznej, pracy oraz kapitału. Ponadto zauważalne są także

tw. dysparytety społeczne i zagrożenia ekologiczne. Można to tłumaczyć na wiele sposobów. Wspomniany wyżej rozdźwięk między sferą makro- i mikroekonomiczną ma swoje uzasadnienie oraz wynika z racjonalnych oczekiwań jednostki, czego nie należy w tym kontekście bezpośrednio przekładać na zasady teorii racjonalnych oczekiwań, które jak wspomniano, nie zostały tu jeszcze omówione. Pojawia się jednak wyraźnie sprecyzowany problem, z którym należy sobie poradzić. Jego istotą jest brak równowagi między wspomnianą skalą makro- i mikroekonomiczną, a czego egzemplifikacją jest nieumiejętność współczesnego człowieka w godzeniu interesu własnego z interesem społeczeństwa jako całości.

2.2. Dylematy oraz koszty przemian w zmiennych uwarunkowaniach niepewności i ryzyka

Zarysowane wyżej dylematy i problemy, którym powinny stawić czoła państwa narodowe o słabnącej funkcji wypełnianej z różnym skutkiem do czasów „przedglobalizacyjnych”, są tym trudniejsze do rozstrzygnięcia i rozwiązania, im bardziej nieadekwatna jest ich międzynarodowa sytuacja w płaszczyźnie społecznej i politycznej wobec wewnętrznych sprzeczności społecznych i gospodarczych, występujących w tych państwach. Jest to związane ze słabnącą rolą państwa w gospodarce poprzez przejęcie zadań dotąd immanentnie państwowych przez rynek, a właściwie korporacje ponadnarodowe. Niezrównoważony rozwój w sferze społecznej i technologicznej jest prostą konsekwencją takiego układu.

Zasygnalizowane problemy niesprzyjające ogólnemu światowemu zrównoważonemu rozwojowi są w istocie silnie uwarunkowane przez zmienność otoczenia społecznego i gospodarczego w różnej skali ekonomicznej. Ich rozwiązanie jest dyktowane charakterem niepewności i ryzyka związanego z poważnymi uwikłaniami wyborów, które z jednej strony mogą prowadzić do degradowania gospodarek państw oraz ich społeczeństw w skali krótkookresowej, co nie jest tu bliżej określone, gdyż może odnosić się do kilkunastu, a nawet kilkudziesięciu lat, a z drugiej strony są one dyktowane celami operacyjnymi i długofalowymi. Generują te procesy problemy i koszty, których oszacowanie ze względu na niepewność i nie do końca wyjaśnione umocowanie instytucjonalne organizacji działających na arenie światowej, nie jest nawet w największym zarysie możliwe. Wszelka prognoza w tym zakresie graniczyłaby z niesamowitą nieodpowiedzialnością, a w znaczeniu naukowym byłaby zupełnie pozbawiona podstaw metodologicznych. Istotne jest zatem to, co powiedziano wcześniej o ocenie racjonalności z punktu widzenia analizy celów.

Zarysowująca się w tym miejscu myśl przewodnia wynika z konstatacji, że ogół jako suma optymalizowanych segmentów nie jest optymalny, gdyż brakuje spójności w optymach segmentów, w znaczeniu gospodarki globalnej i gospodarek państwowych.

Zasygnalizowane powyżej problemy mogą być wyjaśniane w bardzo różnorodny sposób. Pojawiają się bowiem dylematy społeczne, których rozstrzygnięcie w skali makro nie jest możliwe (w najbliższym stuleciu), a wzmaga to nierównomierność wzrostu demograficznego w poszczególnych częściach świata oraz rozproszenie dochodowe. Niezbędne są w tym zakresie szczególne badania procesów ekonomicznych.

Zastanowienie i wyzwanie wzbudza także fakt różnicowania się społecznego aż do przeciwstawności w aspekcie wewnętrznym i wpływ rozwoju gospodarczego na wzrost tej polaryzacji. Zatem jakie działania w przestrzeni globalnej w sensie instytucjonalnym i operacyjnym na płaszczyźnie państw są niezbędne, by stanowisko takie sprzeczne z powszechnie dominującym, a odmiennym od niego, nie stało się faktem? Uzasadnieniem, być może niepełnym, tego stanowiska jest autentyczny proces obserwowany w aspekcie wzrastającej agresji społecznej, terroryzmu światowego przy dominującej sile gospodarek w ujęciu globalnym. Czyżby miały to być właśnie te koszty przemian globalnych?

2.3. Ekonomiczne ryzyko globalne – ryzyko megaekonomiczne

Zasygnalizowany wyżej kontekst postrzegania pewnych współczesnych problemów gospodarczych oraz społecznych wprowadza elementy wyzwolone zmiennością otoczenia mikro- i makroekonomicznego, co prowadzi do ryzyka megaekonomicznego, rozumianego jako nieefektywny pogłębiający się proces odnoszący się do polaryzacji społecznej, gospodarczej i ekologicznej w przestrzeni globalnej. Pokazuje to, że ryzyko megaekonomiczne może wywołać nierównomierności rozwojowe świata, a właściwie może prowadzić do niezrównoważonego rozwoju jeszcze bardziej wyeksponowanego niż ten, z którym mamy do czynienia w obecnym czasie, na początku dwudziestego pierwszego wieku.

Przyjęte stanowisko, przy odwołaniu się do oceny racjonalności z pozycji analizy celów, może prowadzić do dążenia państw w dostosowaniu swych gospodarek poprzez zmianę założeń mikroekonomicznych dla przeciwdziałania obserwowanemu w skali megaekonomicznej ryzyku. Zatem, jak już wspomniano wcześniej, może to potęgować w szerszej skali niekompletność globalizacji, co uzasadniałoby błąd w założeniu tego procesu.

W innym aspekcie można widzieć klasyczne keynesowskie ujęcie w teorii ekonomii, gdy procesy globalizacyjne wzmogą popyt zewnętrzny, co już jest dostrzegalne. Staje się on wtedy substytutem aktywności krajowej, w tym zakresie działalności gospodarczej państwa. Ryzyko megaekonomiczne jest

wtedy asymetrycznie ustawione względem wszelkich działań zmierzających do zarządzania tym ryzykiem w celu jego zminimalizowania. Występują tu jednak oczywiste sprzężenia w działaniach na płaszczyźnie państwa i pewne sprzeczności mogące tak czy inaczej wzmacniać niekompletność globalizacji.

2.4. Rola rynku w aspekcie doskonałej konkurencji w warunkach ryzyka megaekonomicznego

Nie ulega wątpliwości, a przynajmniej jest wyraźnie uzasadnione, że w warunkach przejawiania się ryzyka megaekonomicznego rynek nie ma tych cech, które go charakteryzowały, gdy ryzyko to nie występowało lub było mało aktywne. Wzrost ograniczeń doskonałej konkurencji może być efektem ryzyka megaekonomicznego, a ogólnie samej globalizacji. Podmioty aktywnie uczestniczące na rynku siłą rzeczy dla złagodzenia wzrastających ograniczeń doskonałej konkurencji koncentrują się na strukturach rynkowych poprzez monopolizację i oligopolizację. Prowadzi to do interwencjonizmu państwowego na rzecz podmiotów wyłączonych z tego procesu, które odczuwają dyskomfort wynikający, jak twierdzi się w teorii ekonomii, z redystrybucji wartości dodanej przez mechanizmy rynkowe. Tu jednak w wyniku występowania ryzyka megaekonomicznego rola państwa ma ograniczone możliwości. Jedynym skutecznym możliwym przeciwdziałaniem takim tendencjom może być integrowanie się w szersze struktury tych zagrożonych podmiotów gospodarczych.

2.5. Paradoksy globalizacji i inne kwestie

Nie jest precedensem występowanie obserwowanych obecnie procesów globalizacyjnych w historii gospodarczej świata. Procesy rozwojowe charakteryzowały się zawsze cyklicznością i były oparte na stałej strukturze tych procesów. Jest to schemat wypracowany przez determinizm historyczny (Althusser, Balibar).

Paradoksy w mechanizmach rozwojowych obserwowane obecnie wynikają np. z finansowania przez raczej niezamożne Chiny, które powinny być odbiorcami kapitału, ponadprzeciętnie bogatych Amerykanów. Recesyjna faza cyklu koniunkturalnego może zatrzymać ten proces rozwoju kontaktów na płaszczyźnie ekonomicznej między Chinami a USA.

W obecnie obserwowanym kryzysie jest pewna analogia do okresu sprzed kryzysu z lat trzydziestych ubiegłego wieku. Mikroekonomiczny charakter postrzegania tamtej sytuacji stoi jednak w pewnej sprzeczności z procesami globalizacji w obecnym kształcie, czego efektem nie musi być nieograniczone rozwijanie się procesu kryzysowego, gdyż proces globalizacji nie musi się rozwijać nieuchronnie w sposób ciągły (możliwe są regresy w tym rozwoju). Bezpośrednie przenoszenie sytuacji sprzed osiemdziesięciu lat na obecny kry-

zys wydaje się nieuzasadnione. Istniejące wyraźne granice rozwoju kryzysu wynikają obecnie z behawioralnych zachowań konsumenckich, głównie w krajach rozwiniętych, a w szczególności ze stosunku do społecznej akceptacji nierówności społecznych, nastawienia społeczeństw do zjawisk ekologicznych i innych. Uzasadnienie dla takiej oceny prowadzi do sformułowania wyraźnych *explicite* warunków niwelujących niekompletność procesów globalizacyjnych. Jednymi z nich są instytucjonalne formy ponadnarodowe.

Do paradoksów można zaliczyć przykładowy rozwój budowy sieci autostrad w Chinach w porównaniu ze stanem panującym w Indiach. Zauważalny jest tu styk ekonomii i polityki w kontekście konkurencji w ujęciu globalnym. Wnioski są paradoksalne, gdyż w tej konkurencji odmiennych systemów politycznych ogromnych krajów, wygrywa system autorytarny. Nie są wyjaśnione, podobnie zresztą jak kapitałowe relacje między USA i Chinami, granice autorytarnych działań w sferze gospodarczej. Zarówno w historii starożytnej, jak i dwudziestowiecznej, można by znaleźć pewne analogie tłumaczące takie zjawiska.

2.6. Globalizacja a pozorne potrzeby – kreacja innowacji

Wartość wymienna w procesie globalizacji zdaje się mieć przewagę nad wartością użytkową. Według poglądów niektórych autorów w procesach gospodarczych wzrastające znaczenie ma marketing, z czego wynikają tzw. pozorne potrzeby. W ujęciu teorii ekonomii jest to zaprzeczenie klasycznych zasad A. Smitha, gdyż konsumenci mogą nie zaspokajać potrzeb egzystencjalnych, natomiast rynek nie jest tworzony przez niezależnych konsumentów decydujących o alokacji czynników wytwórczych.

Procesy globalizacyjne są często wyjaśniane przez pryzmat analizy skutków jej mikroekonomicznego charakteru. O dylematach z tego wynikających wspomniano już wcześniej, akcentując moment wstecznego efektu, wynikający z działań mikroekonomicznych na poziomie krajów. Formułuje się w tym zakresie hipotezę nawiązującą do wspomnianej na wstępie irracjonalności globalizacji, że ma ona wpływ na zwiększenie znaczenia marketingu w prowadzonej działalności gospodarczej. Ma to związek z obserwowaną dość powszechnie praktyką, o wpływie właścicieli (w tym środków medialnych) na prezentowane obiegowe i oficjalne opinie. Wydaje się też istotne, że kontekst ten uwzględnia umiejętność dostosowywania się poglądów opinii społecznej lub dystansu do teorii zastanych przez młode pokolenia społeczeństw w wyniku ich wnikliwszego weryfikowania z rzeczywistością. Wniosek jaki z tego wynika sprowadza się do stwierdzenia, że aby kreować nowe innowacje, należy doskonalić wcześniej stworzone.

2.7. Scenariusze dopasowań do modeli globalizacji

Ogólnie rzecz biorąc, można rozróżnić wiele wzorców dostosowań do procesów globalizacji. Spośród nich najbardziej interesujący wydaje się wzorzec: skandynawski, chiński i neoliberalny. Według ocen spotykanych w literaturze model ten ma charakter aktywny, w którym w sposób długookresowy jest przewidziane wsparcie konkurencyjności zasobów oraz kapitału społecznego. W modelu chińskim, bardzo restrykcyjnym z punktu widzenia społecznego, wykorzystuje się ogromne zasoby pracy i wysoką stopę oszczędności. Przewartościowanie kursu walutowego oraz wyłączenie napływu kapitału spekulacyjnego (portfelowego) sprzyja konkurencyjności gospodarki. Ponadto w modelu tym istnieją duże możliwości tworzenia firm przez podmioty zewnętrzne. Model ten cechuje duża rozpiętość pomiędzy rozwojem ekonomicznym a społecznym, co jak na kraj pozornie oparty na modelu społecznym z akcentowaniem władzy ludowej, wydaje się nieskojarzone z rzeczywistością. Natomiast model neoliberalny pojmuje globalizację jako szczybel gospodarki rynkowej ze szczególnym osłabieniem roli państwa, co stanowi priorytet w tym modelu.

2.8. Konkurencja w życiu gospodarczym jako przyczyna internacjonalizacji gospodarczej – badania i ich wymiar

Coraz bardziej intensywne konkurowanie różnych podmiotów ekonomicznych jest wywołane zwiększającym się tempem procesu internacjonalizacji życia gospodarczego w pewnych aspektach mających wymiar globalizacyjny. Proces ten wynika z dążenia narodów, państw i ich związków do osiągnięcia coraz wyższego poziomu dobrobytu. W wyścigu tym uczestniczą także rządy państw i ich ugrupowań, starając się poprzez tworzone ramy instytucjonalno-instrumentalne tworzące system, umożliwić konkurowanie tych podmiotów. Służą temu także wszelkie badania teoretyczne, dotyczące m.in. wzrostu gospodarczego, międzynarodowych stosunków gospodarczych i lokalizacji działalności gospodarczej na świecie. Badania te mają charakter najczęściej deskryptywny, tzn. opisują jak jest. W pewnych przypadkach ich charakter jest normatywny i wtedy umożliwiają udzielenie odpowiedzi, jak powinno być. W pierwszym przypadku badania nie zawsze wyprzedzają rozwój wydarzeń, ale nawet wtedy warto uwzględnić tzw. podstawowe stylizowane fakty z przeszłości, z których większość ma niewiele wspólnego z tzw. doktrynerstwem, z którym walkę sprowadza się czasami nawet do odrzucania podstawowych uniwersalnych praw racjonalnego gospodarowania. Jest tak np. z dynamicznie ujmowaną zasadą przewag względnych w skalach państwowej i między-państwowej, którą interpretuje się obecnie na różne sposoby.

W ogólnym spojrzeniu na badania należy stwierdzić, że rozwój bieżących wydarzeń gospodarczych i ich prognozowanie bez uwzględniania odpowiedniego dorobku teoretycznego, niestety często ograniczonego, nie jest najlepszym rozwiązaniem.

Istnieją w literaturze stwierdzenia, z których wynika, że nie istnieją w istocie żadne prawa ekonomiczne, co odsyła całą wiedzę o gospodarowaniu do swoistego lamusa. Nawet jeśli tego rodzaju stwierdzenia są przesadzone i uczynione dla efektu medialnego, autor takiego stwierdzenia nie wystawia sobie właściwej noty i także takie stwierdzenie staje się nieprawdziwe, by zgadzać się z wyrażonym przez niego jasno „prawem”. Często autorzy przeczą wyrażonym przez siebie stwierdzeniom.

Jeśli w kontekście globalizacji przyjmuje się, że procesy jej towarzyszące oraz uwarunkowania, przebieg i implikacje należy postrzegać w trzech wymiarach: przestrzennym – wymiar geograficzny, czasowym – wymiar historyczny i interdyscyplinarność, to jak stwierdza J. Miś, należy także uwzględnić czwarty wymiar, za A. Marshalllem (ang. *fourth agent of production*). Rozumiał on przez to systemowe uwarunkowania międzynarodowej zdolności konkurencyjnej i bieżącej międzynarodowej konkurencji poszczególnych krajów oraz ich grup. Znaczenie tego czynnika zmieniało się w czasie, ale warto na niego spojrzeć przy odpowiednich rozważaniach teoretycznych i prezentowaniu wyników różnorodnych analiz empirycznych.

2.9. Koncepcja i kontrowersje wokół teorii racjonalnych oczekiwań

Teoria racjonalnych oczekiwań została sformułowana z pozycji ekonomicznych teorii w latach 70., wobec niesprawdzającej się wtedy w problemach makroekonomicznych hipotezy oczekiwań adaptacyjnych. Została ona po raz pierwszy zaproponowana przez amerykańskiego ekonomistę Hojną Mutha w 1961 r. i rozwinięta przez wybitnego ekonomistę Roberta Lucasa, a także Thomasa Sargenta i Neila Wallace’a.

Według teorii racjonalnych oczekiwań podmioty gospodarcze podejmują swoje decyzje na podstawie wszystkich dostępnych informacji o aktualnych uwarunkowaniach ekonomicznych oraz o potencjalnych skutkach tych decyzji. Zakłada się również zdolność podmiotów do formułowania wniosków na podstawie zdarzeń z przeszłości; wtedy są zdolne do antycypacji zdarzeń i zaistnienia możliwych scenariuszy.

Twórcy tej teorii poddali krytyce model Keynesa, zarzucając mu nieprecyzyjność względem oczekiwań podmiotów gospodarczych jako wielkości egzogenicznych spoza gospodarki. Wywoływało to niedoprecyzowanie obrazu

gospodarki, gdzie widoczne są tylko konsekwencje zmian oczekiwań, a nie ich przyczyny. W opinii przedstawicieli szkoły racjonalnych oczekiwań polityka stabilizacyjna oparta na takim modelu prowadziła tylko do wzrostu inflacji.

Teoria racjonalnych oczekiwań podaje w wątpliwość skuteczność angażowania się polityki gospodarczej w dynamizowanie wzrostu gospodarczego, gdyż państwo nie ma wpływu na trwały wzrost zatrudnienia lub produktu. Twórcy tej teorii posuwają się w swoich poglądach jeszcze dalej – uważają, że przeciwdziałanie recesji i bezrobociu poprzez stosowanie aktywnej polityki finansowej, nie przynosi efektów, a pojawienie się bezrobocia nie wynika z niedostatecznego popytu. Według tej teorii państwo powinno dążyć do utrzymania stabilności cen oraz działać w kierunku stabilizowania podaży w gospodarce. Nie zakładają natomiast wbrew keynesistom możliwości stosowania bezpośrednich ingerencji państwa. Stabilizacja cen powinna przebiegać na podstawie wypracowanych reguł, do których podmioty gospodarcze będą miały zaufanie.

Szkół racjonalnych oczekiwań różni się od monetaryzmu innym spojrzeniem na rolę państwa w gospodarce, gdyż dopuszcza jego działania zarówno poprzez oddziaływanie na rynek pieniężny, jak również równoważenie budżetu, sterowanie wydatkami rządowymi i cele realnego kursu walutowego.

Krytycy tej teorii kwestionują racjonalne oczekiwania jako wiarygodny model zachowania przedsiębiorstw, twierdząc na podstawie badań empirycznych, że nie wszystkie podmioty zachowują się w pełni racjonalnie oraz że wiele z nich popełnia cyklicznie te same błędy. Wynika to z faktu, że podmioty na ogół w optymalny sposób wykorzystują tylko część dostępnych informacji. Zwraca się też uwagę na to, że nawet przy w pełni racjonalnym zachowaniu istnieje miejsce dla aktywnej roli polityki rządowej, gdy cenowe i popytowe mechanizmy dostosowawcze nie zadziałają na czas.

Należy wspomnieć, że teoria racjonalnych oczekiwań została na początku lat 80. poddana totalnej krytyce przez naukowca polskiego pochodzenia Richarda Frydmana i Edmunda Phelps, lecz ostracyzm środowiska naukowców z Chicago University spowodował zmianę stanowiska tego ostatniego, który dzięki temu w 2006 r. został noblistą, a Frydman, który trwał przy swym stanowisku, został praktycznie wykluczony ze środowiska. Dopiero w 2007 r. wraz z Michaelem Goldbergem wydał znaczącą pracę *Imperfect Knowledge Economics. Exchange Rates and Risk*, która znalazła uznanie także u E. Phelps.

MODEL CONTROVERSIAL AXIOLOGICAL CONTEXT OF RATIONAL EXPECTATIONS IN ECONOMICS

Summary

The article presents selected aspects of the controversial axiology of rational expectations in the background generally recognized perception of economic models and their exemplification. Uncertainty and volatility risk in the dimension of economic processes are presented in general presented the processes of globalization, competition and innovation. The role of the market is mentioned in the context of perfect competition in terms of economic mega risks. Introductory considerations are in any kind of detailed statistical and econometric terms of economic and financial matters dealt with in formal terms, a model. The need for a holistic view of the ongoing social and economic processes is prompted research into the condition of the economy on a global basis, taking into account regional differences.

UOGÓLNIONA METODA TARYFIKACYJNA KLASY MBP W UBEZPIECZENIACH KOMUNIKACYJNYCH

Wprowadzenie

Zakłady ubezpieczeń działające w grupie ubezpieczeń majątkowych tworzą własne systemy taryfikacyjne w celu ustalenia sprawiedliwego poziomu składki dla każdego ryzyka w różnego rodzaju portfelach ubezpieczeniowych. W artykule omawiany jest proces tzw. taryfikacji *a priori* polegający na kalkulacji składki początkowej, biorąc pod uwagę charakterystyczne cechy opisujące osobę ubezpieczającą się (wiek, płeć itd.) oraz przedmiot ubezpieczenia (pojazd, nieruchomość itd.). Cechy te wyrażane są poprzez zmienne taryfikacyjne (zero-jedynkowe), natomiast ich wpływ na wielkości szkód w portfelu mierzą stopy taryfy (ang. *relativities*). Realizację zmiennych taryfikacyjnych determinują podział portfela na grupy taryfikacyjne. Po wyznaczeniu składki początkowej w każdej grupie oblicza się składkę dodatkową (liczoną najczęściej jako procent składki początkowej), gdzie uwzględnia się liczbę szkód spowodowanych w przeszłości przez osobę ubezpieczającą się, stosując najczęściej system *bonus – malus* (system zniżek – wyżek). Istotnym problemem w zakładzie ubezpieczeń jest zatem ustalenie zmiennych taryfikacyjnych. Obserwuje się zasadniczo wzrost liczby czynników uwzględnianych przez zakłady w kalkulacji składki, co jest przedstawione krótko w następnym rozdziale.

W praktyce w zakładach ubezpieczeń najczęstsze zastosowanie w budowie systemów taryfikacyjnych mają uogólnione modele liniowe (GLM). Jednak często w celu lepszego zrozumienia natury danych, warto jest sięgnąć we wstępnej analizie do prostszych metod taryfikacji, tzw. procedur minimalnego obciążenia (ang. *Minimum Bias Procedures – MBP*). To podejście pierwszy raz zaprezentowano w pionierskiej pracy Bailey’a oraz LeRoy’a¹. Statystyczne metody przedstawione w tej pracy pozwalają na szacowanie poziomu zmiennych taryfikacyjnych przy minimalnym poziomie funkcji

¹ R.A. Bailey, J.S. LeRoy, *Two Studies in Automobile Insurance Ratemaking*, PCAS XLVII, 1960, s. 1-19.

obciążenia (ang. *bias function*). W zależności od wyboru postaci modelu (addytywny czy multiplikatywny) oraz funkcji obciążenia uzyskiwane są stopy taryf klasyfikujące ryzyka poprzez podział portfela niejednorodnego na sub-portfele jednorodne. W literaturze przedmiotu można znaleźć zarówno liczne rozwinięcia klasycznych czterech metod MBP, jak i powiązanie podejścia minimalizacji funkcji obciążenia z modelami GLM².

W niniejszym artykule przedstawiono tzw. uogólnioną metodę MBP, która pozwala w łatwy sposób dokonać szybkiej analizy wstępnej podziału portfela jednorodnego na grupy taryfikacyjne poprzez zastosowanie algorytmu iteracyjnego wyznaczania wskaźników wpływu dla zmiennych taryfikacyjnych przy różnych kombinacjach parametrów. Działanie metody oraz algorytmu przedstawiono na przykładzie empirycznym, wykorzystując dane zaczerpnięte z literatury. Obliczenia wykonano w programie komputerowym R.

1. Krótka charakterystyka taryfikacji na polskim rynku ubezpieczeń komunikacyjnych

W celu przybliżenia możliwości wykorzystywania prezentowanych w artykule metod taryfikacji, poniżej przedstawiono krótko praktykę polskiego rynku ubezpieczeniowego na przykładzie PZU S.A. Jeżeli chodzi o polski rynek ubezpieczeń komunikacyjnych, garść informacji dotyczących taryfikacji dostarcza Wspólny Raport Urzędu Komisji Nadzoru Finansowego i Ubezpieczeniowego Funduszu Gwarancyjnego „Ubezpieczenia komunikacyjne w latach 2005-2009”. Raport wskazuje, iż najczęściej zakłady ubezpieczeń uzależniały wysokość składki od: marki pojazdu, rodzaju pojazdu, pojemności silnika, miejsca zamieszkania. Dodatkowe czynniki to: charakter użytkowania pojazdu, wiek posiadacza pojazdu, okres posiadania prawa jazdy, model pojazdu, rodzaj płatności, posiadanie innych umów w zakładzie ubezpieczeń. Niektóre zakłady ubezpieczeń wprowadzają dodatkowe zmienne taryfikacyjne: płeć, posiadanie dzieci, przynależność do określonej grupy zawodowej, rodzaj paliwa, liczba użytkowników pojazdu, liczba ubezpieczonych pojazdów, ważność badań technicznych. Analiza zasad ustalania przez zakłady ubezpieczeń składki w ubezpieczeniu OC posiadaczy pojazdów mechanicznych pokazuje, iż zakłady zwiększają ilość zmiennych taryfikacyjnych wpływających na ostateczny poziom składki. Można ponadto zauważyć, iż ocena ryzyka zależy od większej liczby czynników w zakładach ubezpieczeń działających w systemie „direct”.

² S. Mildenhall: *A Systematic Relationship between Minimum Bias and Generalized Linear Models*. PCAS, LXXXVI, 1999, s. 394-487.

Przykładowe zmienne taryfikacyjne stosowane przez PZU S.A. (ubezpieczenie OC dla posiadaczy pojazdów mechanicznych TI) to: klasa taryfowa według pojemności silnika, strefa regionalna (składka bazowa), wiek posiadacza pojazdu, prawo jazdy krócej niż 2 lata, klasa bonus-malus, okres eksploatacji pojazdu (% składki podstawowej). Poniższe zestawienie przedstawia czynniki (zmienne taryfikacyjne) uwzględniane w systemie taryfikacyjnym w zakładzie ubezpieczeń PZU S.A. (na podstawie sposobu naliczania składki w ubezpieczeniu OC posiadaczy pojazdów mechanicznych TI).

Tabela 1

Czynniki taryfikacyjne z PZU S.A.

Czynniki	Opis
klasa taryfowa	klasa taryfowa ustalana jest według pojemności silnika
strefa regionalna	zniżka lub wyżka
klasa pojazdu	
ubezpieczenie krótkoterminowe	
wiek posiadacza pojazdu	dotyczy, jeżeli pojazd zarejestrowany jest czasowo
klasa bonus-malus	poniżej 23 lat, poniżej 23 lat i 12 miesięcy bezszkodowych, powyżej 22 i poniżej 26 lat (w innym przypadku wiek nie wpływa na składkę)
prawo jazdy krócej niż 2 lata	dotyczy posiadacza pojazdu
okres eksploatacji pojazdu	poniżej 1 roku, powyżej 1 roku i poniżej 2 lat
kontynuacja ubezpieczenia w PZU S.A.	nie dotyczy ubezpieczenia krótkoterminowego
pojazd w leasingu	
jednorazowa opłata składki	

Źródło: TARYFA SKŁADEK za obowiązkowe ubezpieczenie odpowiedzialności cywilnej posiadaczy pojazdów mechanicznych dla klienta indywidualnego oraz małego i średniego podmiotu gospodarczego ustalona Uchwałą Zarządu PZU S.A. Nr UZ/522/2005 z dnia 3 listopada 2005.

W systemie taryfikacyjnym PZU S.A. składka bazowa jest uzależniona jedynie od klasy taryfowej i strefy regionalnej. Oznacza to, iż w taryfikacji *a priori* wprowadzono dwie zmienne taryfikacyjne. Drugim etapem taryfikacji *a priori* jest kalkulacja składki dla wszystkich grup ryzyka wyodrębnionych w procesie doboru zmiennych taryfikacyjnych.

2. Uogólniona procedura MBP wyznaczania wskaźników wpływu

Klasycznie w procedurach MBP występują cztery podstawowe metody szacowania wskaźników wpływu: zasada bilansu, minimalizacja średniego błędu kwadratowego, minimalizacja błędu χ^2 oraz maksymalizacja funkcji wiarygodności³. W ostatnim przypadku niezbędne jest przyjęcie założenia o postaci rozkładu średniej wartości szkody. W literaturze przedmiotu przeprowadzono wiele analiz, które pokazały, iż rozkłady szkód w ubezpieczeniach charakteryzują się zasadniczo następującymi własnościami: przyjmują wartości dodatnie, występuje proporcjonalność wariancji do wartości oczekiwanej oraz asymetria dodatnia⁴. Dlatego też w metodzie maksymalizacji funkcji wiarygodności najczęściej szacuje się wskaźniki wpływu przy założeniu rozkładu gamma lub rozkładu odwróconego normalnego dla szkód. Można również łatwo wykazać, że zasada bilansu jest równoważna z maksymalizacją funkcji wiarygodności przy założeniu rozkładu Poissona.

Uogólniona metoda MBP jest rozszerzeniem podejścia, w którym maksymalizuje się funkcję wiarygodności, gdzie wskaźniki wpływu szacuje się w zależności od dwóch parametrów: p oraz q . Dla poszczególnych kombinacji p i q uzyskiwane są estymatory wskaźników wpływu analogicznie do przypadku maksymalizacji funkcji wiarygodności dla różnych rozkładów:

- a) $p = 1, q = 1$ – rozkład Poissona,
- b) $p = 1, q = 0$ – rozkład gamma,
- c) $p = 1, q = -1$ – rozkład odwrócony normalny.

Ponadto dla kombinacji parametrów $p = 1, q = 2$ uzyskuje się estymatory analogiczne do tych z metody minimalizacji średniego błędu kwadratowego.

W celu szczegółowej prezentacji omawianej procedury przyjmijmy następujące oznaczenia:

- 1. X_1, X_2 – zmienne taryfikacyjne, Y – zmienna losowa opisująca średnią wartość szkody.
- 2. $x_{11}, \dots, x_{1n}$ – wskaźniki wpływu dla zmiennej X_1 , $x_{21}, \dots, x_{2m}$ – wskaźniki wpływu dla zmiennej X_2 .

³ R.A. Bailey, J.S. LeRoy, op. cit.

⁴ P. McCullagh, J.A. Nelder, *Generalized Linear Models*, Chapman & Hall, New York 1999.

3. y_{ij} – realizacja zmiennej Y (średnia wartość pojedynczej szkody) dla i -tej realizacji zmiennej X_1 oraz dla j -tej realizacji zmiennej X_2 .
4. n_{ij} – liczba szkód dla i -tej realizacji zmiennej X_1 oraz dla j -tej realizacji zmiennej X_2 .
5. B – bazowa wartość szkody.

Zakładając model multiplikatywny dla zmiennych taryfikacyjnych oraz dwie zmienne taryfikacyjne, szacowana średnia wartość szkody dana jest wzorem:

$$\hat{y}_{ij} = Bx_{1i}x_{2j}, \quad i = 1, \dots, n, \quad j = 1, \dots, m$$

Estymatory dla wskaźników wpływu uzyskane, maksymalizując funkcję wiarygodności, można przedstawić w ogólnej postaci w zależności od parametrów p, q ⁵:

$$\hat{x}_{1i}^{s+1} = \frac{\sum_j n_{ij}^p y_{ij} (Bx_{2j}^s)^{q-1}}{\sum_j y_{ij}^p (x_{2j}^s)^q}$$

$$\hat{x}_{2j}^{s+1} = \frac{\sum_i n_{ij}^p y_{ij} (Bx_{1i}^s)^{q-1}}{\sum_i y_{ij}^p (x_{1i}^s)^q}$$

W powyższych dwóch wzorach znacznik s oznacza numer iteracji, gdyż wyznaczenie wartości wskaźników wymaga zastosowania algorytmu iteracyjnego ze względu na wzajemną zależność. Proponowany algorytm przedstawia się w kilku krokach:

KROK 1. Przyjęcie wartości wejściowych dla $s = 1$

$$x_{1i}^1 = 1, \quad x_{2j}^1 = 1$$

KROK 2. Przyjęcie wartości bazowej B

KROK 3. Wyznaczenie szacunkowych wartości $\hat{x}_{1i}$, $\hat{x}_{2j}$

KROK 4. Wyznaczenie szacunkowej wartości $\hat{y}_{ij}$

KROK 5. Wyznaczenie błędu oszacowania

⁵ L. Fu, Ch.P. Wu, *General Iteration Algorithms for Classification Ratemaking*, „Variance Journal” 2007, Vol. 1, Iss. 2.

Jako błąd oszacowania przyjęto procentowe ważone odchylenie przeciętne

$$wapb = \frac{\sum_{i,j} n_{ij} \frac{|y_{ij} - Bx_{1i}x_{2j}|}{Bx_{1i}x_{2j}}}{\sum_{i,j} n_{ij}}.$$

W wyborze ostatecznej postaci modelu przyjęto kryterium minimalizacji kryterium $wapb$. Warunkiem końcowym algorytmu iteracyjnego jest:

$$x_{1i} = \lim_{k \rightarrow \infty} \hat{x}_{1i}^k, \quad x_{2j} = \lim_{k \rightarrow \infty} \hat{x}_{2j}^k.$$

W powyższym algorytmie kryterium startu dla wskaźników wpływu przyjęto na poziomie 1, co oznacza uzyskanie w pierwszym kroku szacunkowej wartości $\hat{y}_{ij}$ równej przyjmowanej wartości bazowej. Jako wartość bazową można przyjąć średnią ważoną wartość szkody dla portfela, traktując ją jako składkę kolektywną. Inne alternatywne wartości bazowe, jak np. składka wiarygodna, dają znacznie gorsze rezultaty, co zostało przetestowane przez autora dla różnych kombinacji p, q .

3. Uogólniona procedura MBP dla empirycznych danych komunikacyjnych

Poniżej przedstawiono przykład obliczeniowy obrazujący działanie uogólnionej procedury MBP. W przykładzie analizowano wartości szkód y_{ij} i przyjęto, iż wpływają na nie dwie zmienne taryfikacyjne z następującymi realizacjami:

- a) wiek kierowcy X_{1j} , $j = 1, \dots, 8$ – wyodrębniono 8 przedziałów wiekowych (patrz tabele poniżej),
- b) sposób użytkowania pojazdu X_{2i} , $i = 1, \dots, 4$ – wyodrębniono 4 sposoby:
 - X_{21} – prywatnie,
 - X_{22} – do pracy poniżej 10 km,
 - X_{23} – do pracy powyżej 10 km oraz 10 km,
 - X_{24} – służbowo.

Ponadto zostały wykorzystane dane dotyczące liczby szkód n_{ij} .

W przykładzie przyjęto następujące założenia:

- a) szkody w portfelu są opisywane poprzez średnią wartość pojedynczej szkody oraz liczbę szkód,

- b) pomiędzy zmiennymi taryfikacyjnymi występuje zależność multiplikatywna,
- c) wartość bazowa – średnia ważona wartość szkody dla portfela.

Algorytm iteracyjny został zaimplementowany w języku R. Dane zaczerpnięto z literatury przedmiotu ze względu na duże trudności w pozyskiwaniu danych od ubezpieczycieli funkcjonujących na polskim rynku. Warunek stopu przyjęto jako $\varepsilon < 0,0000001$. W celu porównywania uzyskanych estymatorów wskaźników wpływu w zależności od parametrów p, q przeprowadzono losowanie tych parametrów z rozkładu jednostajnego. Parametr p losowano z przedziału $[0,2]$, natomiast parametr q z przedziału $[-2,2]$. Teoretycznie możliwe jest przyjmowanie dowolnych wartości przez parametry, jednak wyniki poza tymi przedziałami nie dawały zadowalających rezultatów.

Tabela 2

Wartości szkód oraz liczby szkód w poszczególnych kategoriach (w j.p.)

y_{ij}	17-20 X_{11}	21-24 X_{12}	25-29 X_{13}	30-34 X_{14}	35-39 X_{15}	40-49 X_{16}	50-59 X_{17}	60+ X_{18}
X_{21}	250,48	213,71	250,57	229,09	153,62	208,59	207,57	192
X_{22}	274,78	298,6	248,56	228,48	201,67	202,8	202,67	196,33
X_{23}	244,52	298,13	297,9	293,87	238,21	236,06	253,63	259,79
X_{24}	797,8	362,23	342,31	367,46	256,21	352,49	340,56	342,58
n_{ij}	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
X_{21}	21	63	140	123	151	245	266	260
X_{22}	40	171	343	448	479	970	859	578
X_{23}	23	92	318	361	381	719	504	312
X_{24}	5	44	129	169	166	304	162	96

Źródło: S. Mildenhall: *A Systematic Relationship between Minimum Bias and Generalized Linear Models*, PCAS, LXXXVI, s. 394-487, 1999.

Średnia wartość szkody ważona liczbą polis w całym portfelu wynosi 241,46 j.p. W podziale na wiek oraz sposób użytkowania pojazdu wartość ta kształtuje się odpowiednio: 255,43 j.p. oraz 258,91 j.p. Po przeprowadzeniu algorytmu iteracyjnego (Aneks) uzyskano następujące poziomy stóp taryf w poszczególnych modelach w zależności od przyjętych parametrów p, q .

Tabela 3

Stopy taryf dla modelu $p = 1, q = 1$

Stopa taryfy	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
x_i	1,9810	1,9222	1,7863	1,7281	1,3800	1,5083	1,5294	1,5014

Stopa taryfy	Prywatnie	Do pracy <10km	Do pracy >10km	Służbowo
y_j	0,8507	0,8863	1,0737	1,3965

Źródło: Obliczenia własne.

Tabela 4

Stopy taryf dla modelu $p = 1, q = 0$

Stopa taryfy	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
x_i	1,9500	1,9408	1,7992	1,7241	1,3883	1,5019	1,5249	1,4918

Stopa taryfy	Prywatnie	Do pracy <10km	Do pracy >10km	Służbowo
y_j	0,8509	0,8865	1,0755	1,3990

Źródło: Obliczenia własne.

Tabela 5

Stopy taryf dla modelu $p = 1, q = -1$

Stopa taryfy	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
x_i	1,9328	1,9559	1,8102	1,7203	1,3939	1,4982	1,5217	1,4838

Stopa taryfy	Prywatnie	Do pracy <10km	Do pracy >10km	Służbowo
y_j	0,8509	0,8867	1,0771	1,4017

Źródło: Obliczenia własne.

Tabela 6

Stopy taryf dla modelu $p = 1, q = 2$

Stopa taryfy	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
x_i	2,0311	1,9008	1,7718	1,7321	1,3689	1,5178	1,5354	1,5128

Stopa taryfy	Prywatnie	Do pracy <10km	Do pracy >10km	Służbowo
y_j	0,8500	0,8857	1,0712	1,3948

Źródło: Obliczenia własne.

W modelu pierwszym błąd wyniósł $wapb = 0,0445$, w modelu drugim $wapb = 0,0421$, w modelu trzecim $wapb = 0,0415$, natomiast w modelu czwartym $wapb = 0,047$. Najlepszy poziom oszacowania uzyskano zatem w modelu trzecim, gdzie parametry $p = 1$ oraz $q = -1$. Tabela taryfikacyjna dla tego modelu ma następującą postać.

Tabela 7

Tabela taryfikacyjna dla modelu $p = 1, q = -1$

	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
Prywatnie	1,6447	1,6643	1,5403	1,4638	1,1861	1,2749	1,2948	1,2626
Do pracy <10km	1,7139	1,7344	1,6051	1,5254	1,2360	1,3285	1,3493	1,3157
Do pracy >10km	2,0819	2,1068	1,9498	1,8529	1,5014	1,6138	1,6391	1,5983
Służbowo	2,7092	2,7416	2,5372	2,4112	1,9537	2,1000	2,1329	2,0798

Źródło: Obliczenia własne.

Przyjmując składkę bazową na poziomie średniej ważonej szkody w całym portfelu 241,46, szacowane wartości szkód w poszczególnych grupach taryfikacyjnych przedstawia tabela 8.

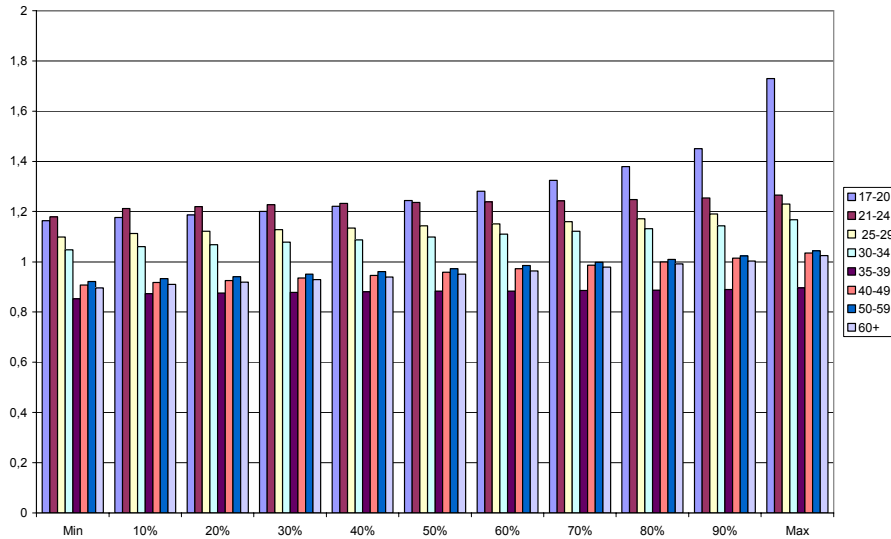
Tabela 8

Szacowane wartości szkód w grupach taryfikacyjnych

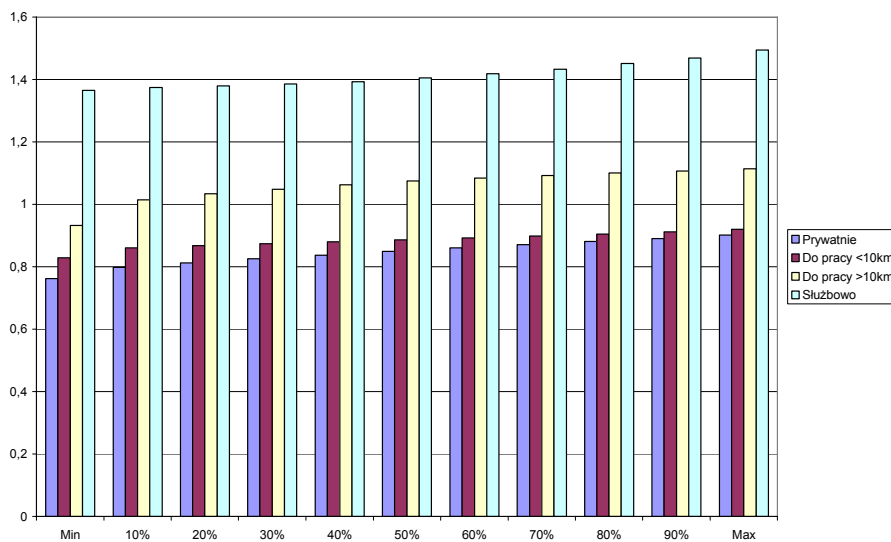
	17-20	21-24	25-29	30-34	35-39	40-49	50-59	60+
Prywatnie	397	402	372	353	286	308	313	305
Do pracy <10km	414	419	388	368	298	321	326	318
Do pracy >10km	503	509	471	447	363	390	396	386
Służbowo	654	662	613	582	472	507	515	502

Źródło: Obliczenia własne.

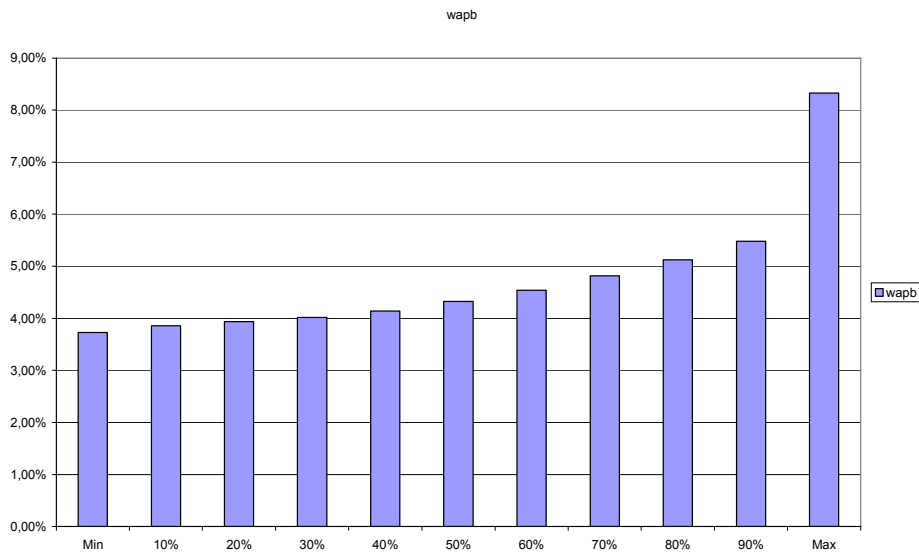
Dla wybranego modelu przeprowadzono również symulację Monte Carlo, przyjmując dla wartości szkody rozkład gamma. Poniższe rysunki przedstawiają rozkłady percentylowe wskaźników wpływu oraz błędu *wapb* w poszczególnych grupach taryfikacyjnych.



Rys. 1. Rozkłady wskaźników wpływu dla wieku



Rys. 2. Rozkłady wskaźników wpływu dla sposobu użytkowania pojazdu



Rys. 3. Rozkład parametrów błędów *wapb*

Rysunki te pokazują, jak mogą się zmieniać stopy taryf w zależności od zmian poziomu szkodowości w poszczególnych grupach taryfikacyjnych. Widać większe zróżnicowanie w grupie – użytkowanie pojazdu w stosunku do grupy wiekowej.

Podsumowanie

W pracy przedstawiono uogólnioną metodę z klasy MBP, która rozszerza podejście maksymalizacji funkcji wiarygodności. W celu prezentacji działania procedury zaimplementowano algorytm iteracyjny szacowania wskaźników wpływu w programie R. W praktyce konieczność implementacji komputerowej algorytmu dla metod MPB jest postrzegana jako główna wada tego typu podejścia do taryfikacji *a priori*. Podejście uogólnione pozwala jednak na szybkie przeprowadzenie taryfikacji dla różnych kombinacji parametrów p oraz q . Ponadto wydaje się, iż ogólna idea metod klasy MBP jest zdecydowanie bardziej intuicyjna niż dość zawiły aparat modeli GLM, co powoduje możliwość lepszego zrozumienia uzyskiwanych wyników i późniejszego ostatecznego wyboru modelu taryfikacyjnego.

GENERALIZED MBP RATEMAKING IN COMMUNICATIONS INSURANCE**Summary**

The paper discusses the a priori ratemaking, which is used to build systems of tariffs for non-life insurance companies. In practice, most companies use generalized linear models GLM. But often in order to better understand the nature of the data, it is useful to reach the preliminary analysis of the simpler methods of ratemaking called the minimum bias procedures MBP to allow for estimating the level of tarification variables by minimizing the bias function. The paper presents the generalized method of MBP, which allows to easily make a quick analysis of pre-classifying the portfolio into groups of uniform tarification. Uses an iterative algorithm for determining indicators of the impact of tarification variables with different combinations of parameters, which are implemented in a computer program R.

Aneks

Implementacja uogólnionej metody klasy MBP w programie R

```

p=runif(1,-2,2) #parametry metody
q=runif(1,-2,2) #parametry metody
p_sim=NULL
q_sim=NULL
x_sim=NULL # wskaźniki wpływu dla zmiennej wiek
y_sim=NULL # wskaźniki wpływu dla zmiennej sposób użytkowania
pojazdu
for ( sim in 1:1000){
y0=c(1,1,1,1)
x0=c(1,1,1,1,1,1,1,1)
x=NULL
y=NULL
B=c()
B=c(B, sum(dane*wagi)/sum(wagi))
macierz.Bxy=NULL
macierz.Bxy=matrix(nrow=nrow(dane), ncol=ncol(dane))
x=rbind(x, x0)
y=rbind(y, y0)
for (s in 1:10) {
  x_temp=c()
  for (m in 1:ncol(dane)){

x_temp[m]=sum(dane[,m]*(wagi[,m]^p)*((y[s,]*B[1])^(q-
1)))/sum((wagi[, m]^p)*((y[s,]*B[1])^q))
  x=rbind(x, x_temp)
  y_temp=c()
  for (n in 1:nrow(dane)){
    y_temp[n]=sum(dane[n,]*(wagi[n,]^p)*((x[s+1,]*B[1])^(q-
1)))/sum((wagi[n,]^p)*((x[s+1,]*B[1])^q))
  y=rbind(y, y_temp)
for (i in 1:nrow(dane)) {
  for (j in 1:ncol(dane)) {

```

```
        macierz.Bxy[i,j]=B[1]*x[nrow(x),j]*y[nrow(y),i]}}  
B=c(B,sum(macierz.Bxy*wagi)/sum(wagi))}      #oszacowane    wartości  
szkody
```

Monika Dyduch

WYCENA PRODUKTU STRUKTURYZOWANEGO NA PRZYKŁADZIE LOKATY STRUKTURYZOWANEJ STABILNA ZŁOTÓWKA

Wprowadzenie

Wahania cen akcji i niepewna przyszłość inwestycji na rynku kapitałowym skłoniły wielu inwestorów indywidualnych do poszukiwania alternatywnych i bezpiecznych form inwestowania. Instytucje finansowe szybko odpowiedziały na potrzeby rynku, oferując szeroką gamę lokat strukturyzowanych, które posiadają zarówno dobre, jak i złe strony. Do najważniejszych zalet, dla których warto inwestować w instrumenty strukturyzowane notowane na giełdzie, można zaliczyć:

- możliwość ochrony kapitału przy jednoczesnym udziale w zyskach (produkty gwarantujące ochronę kapitału),
- możliwość osiągania zysków przewyższających tradycyjne formy oszczędzania, takie jak lokaty i obligacje,
- znaną formułą wypłaty, którą w momencie wykupu oferuje dany instrument (wypłata jest określona formułą znaną inwestorowi przez rozpoczęciem inwestycji),
- możliwość dostępu do nowych rynków (np. zagranicznych) i nowych instrumentów (m.in. surowców, indeksów, walut), do których dostęp był dotychczas niemożliwy – szczególnie w przypadku wielu inwestorów indywidualnych,
- możliwość bezpośredniego dostępu do usług, które do tej pory były niedostępne dla inwestorów indywidualnych,
- w przypadku produktów strukturyzowanych notowanych na giełdzie możliwość sprzedaży instrumentu na rynku wtórnym (inwestor ma możliwość wycofania się z inwestycji przed terminem wykupu instrumentu przez emitenta),
- dywersyfikację inwestycji.

Przedstawione zalety budzą zainteresowanie również z perspektywy płaszczyzny naukowej, tym bardziej, że polski rynek produktów strukturyzowanych jest rynkiem młodym i pewne zachowania, zarówno emitentów, jak i klientów, dopiero się kształtują. Emitenci muszą częściej dostosowywać ofertę do możliwości oraz wiedzy klientów, a klienci z kolei powinni nauczyć się akceptować inwestycyjne porażki, ponieważ produkty strukturyzowane nie są uniwersalnym sposobem na zarabianie pieniędzy, mimo że ich zadaniem jest uchronienie inwestora przed stratami, a przy tym danie mu szansy na dodatkowe zyski przy realizacji określonego scenariusza rynkowego¹.

W pracy podjęto próbę wyceny produktu strukturyzowanego na przykładzie lokaty strukturyzowanej emitowanej przez mBank, której indeksem podstawowym jest kurs fixingu EUR/PLN ogłaszany przez NBP. Wycena produktu sprowadza się w tym przypadku do wyceny instrumentu bazowego, czyli do oszacowania przyszłej wartości indeksu podstawowego. Prognoza instrumentu bazowego jest wygenerowana na podstawie modelu skalającego analizę falkową i sieci neuronowe.

Model zastosowany do badania opiera się głównie na sieci neuronowej, z uwagi na fakt, że posiadają one kilka interesujących dla badaczy cech, m.in.²:

1. Nie wymagają programowania.
2. Posiadają zdolność uogólnienia.
3. Inteligentne zachowanie, które można określić jako:
 - autoasocjacja – struktury są przedstawiane wielokrotnie, a system ma je zapamiętać i przypomnieć sobie, gdy przedstawi się im podobne, tj. skojarzyć,
 - asocjacja struktur – struktury są przedstawiane wielokrotnie, parami (x,y), pojawienie się x ma wywołać y,
 - pamięć adresowalna kontekstowo – wydobywanie informacji, nie przez znajomość miejsc, ale i atrybutów informacji,
 - klasyfikacja (diagnoza, rozpoznanie) – przypisanie do jednej z ustalonych kategorii,
 - detektor regularności – struktury pojawiają się z pewnym prawdopodobieństwem i należy wykryć statystyczne regularności, tworząc nowe kategorie,
 - optymalne spełnianie ograniczeń – wiele hipotez, które na razie nie mogą być prawdziwe, szukanie kompromisu.

¹ <http://www.structus.pl/sites/default/files/Structus>.

² J. Stefanowski, K. Krawiec, wykład, Instytut Informatyki Politechniki Poznańskiej, Poznań 1995.

4. Równoległe przetwarzanie informacji (w przypadku implementacji sprzętowych).
5. Łatwość użycia – sieci neuronowe w praktyce same konstruuja potrzebne użytkownikowi modele, ponieważ automatycznie *uczą się na podanych przez niego przykładach*. Odbywa się to w taki sposób, że użytkownik sieci gromadzi reprezentatywne dane pokazujące jak manifestuje się interesująca go zależność, a następnie uruchamia *algorytm uczenia*, który ma na celu automatyczne wytworzenie w pamięci sieci potrzebnej struktury danych. Opierając się na tej samodzielnie stworzonej strukturze danych, sieć realizuje potem wszystkie funkcje związane z eksploatacją utworzonego modelu. Chociaż użytkownik potrzebuje pewnej, w głównej mierze empirycznej wiedzy dotyczącej sposobu wyboru i przygotowania danych stanowiących przykłady, a także wyboru właściwego rodzaju sieci neuronowej oraz sposobu interpretacji rezultatów, to jednak poziom wymaganej od użytkownika wiedzy teoretycznej, niezbędnej do skutecznego zbudowania modelu przy stosowaniu sieci neuronowych, jest znacznie niższy niż w przypadku stosowania tradycyjnych metod statystycznych³.
6. Rozproszony charakter przetwarzania informacji i wynikająca z niego:
 - odporność na uszkodzenie/eliminację znacznej nawet części neuronów,
 - odporność na uszkodzone i zaszumione wzorce.

Sieci neuronowe posiadają również wiele wad, wśród których można wymienić:

- powolność większości algorytmów uczących,
- trudności z interpretacją wiedzy nabytej przez sieć (brak lub słabe własności eksplikatywne) w związku z jej (tj. wiedzy) rozproszeniem w sieci (tzw. *distributed knowledge representation*); *czarna skrzynka (blackbox)*,
- trudności z reprezentacją niektórych typów danych, np. cech/atributów nominalnych o wartościach niepodlegających uporządkowaniu; konieczność stosowania kodowania „*1 of n*”,
- duża ilość parametrów (sieci i algorytmu uczącego), przy jednoczesnym braku ścisłych reguł do estymacji ich wartości.

³ <http://www.statsoft.pl/textbook/stathome.html>.

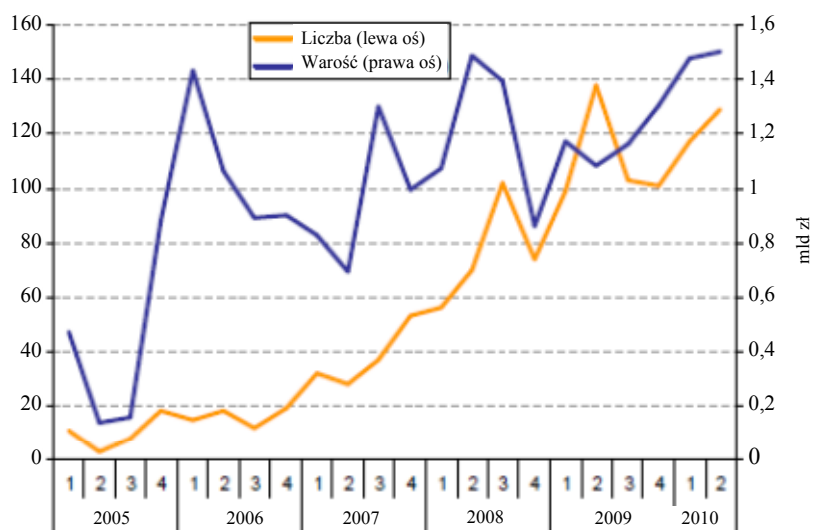
1. Produkty strukturyzowane

Inżynierowie finansowi banku, projektując produkt strukturyzowany, współpracują z bankami inwestycyjnymi, od których kupują odpowiednie instrumenty finansowe (opcje pozwalające zarabiać na rozmaitych rynkach oraz obligacje gwarantujące zwrot części kapitału). Do podstawowych elementów, które ma określony każdy produkt ustrukturyzowany, zalicza się:

- forma prawna (opakowanie) – mimo ogólnej stosowanej nazwy produktu strukturyzowanego, instrument ten może przybierać różne formy prawne, tzw. opakowania. Każda forma niesie różne korzyści, ale również ograniczenia dla inwestora danego produktu. Większość polskich produktów strukturyzowanych ma formę jedną z charakterystycznych strategii inwestycyjnych, tzn.: obligacji strukturyzowanej, lokaty strukturyzowanej, polisy strukturyzowanej, certyfikatu strukturyzowanego, funduszu strukturyzowanego i certyfikatu inwestycyjnego;
- okres zapadalności;
- instrument bazowy – instrument, na którym opiera się zysk, np.: akcje, indeksy, waluty, surowce, fundusze inwestycyjne, stopy procentowe, ratingi (spółek, państw), poziom inflacji;
- konstrukcja wypłaty zysku;
- partycypacja – określony procentowo udział inwestora w zmianie ceny instrumentu finansowego, na którym „zarabia” produkt ustrukturyzowany; zależności współczynnika: $<100\%$ – inwestor zarabia wolniej od instrumentu; $=100\%$ – inwestor zarabia równolegle z instrumentem; $>100\%$ – inwestor zarabia szybciej niż instrument;
- ochrona kapitału.

W ostatnich latach nastąpił systematyczny wzrost liczby oferowanych produktów strukturyzowanych i trend ten nie został odwrócony w czasie ostatniego kryzysu finansowego i gospodarczego⁴. Szczegółowe tendencje w zakresie liczby i wartości sprzedanych produktów strukturyzowanych prezentuje rysunek 1.

⁴ Urząd Komisji Nadzoru Finansowego, Warszawa, październik 2010.



Rys. 1. Liczba i wartość sprzedanych produktów strukturyzowanych w Polsce (dane kwartalne)

Źródło: www.structuredretailproducts.com.

Okres dynamicznego rozwoju rynku produktów strukturyzowanych w Polsce przypada na okres po roku 2006. Od tego czasu przyrasta wolumen sprzedaży, a począwszy od roku następnego – również liczba oferowanych produktów strukturyzowanych. W 2008 r. wielkość sprzedaży produktów strukturyzowanych osiągnęła poziom 4,8 mld zł. W następnym roku wzrost był dużo mniejszy. Wartość produktów strukturyzowanych znajdujących się w portfelach klientów na koniec pierwszego półrocza 2010 r. osiągnęła wartość 15,5 mld zł⁵.

Rynek produktów strukturyzowanych w Polsce nadal nie jest znaczący na tle innych krajów, zarówno w ujęciu bezwzględnym, jak i w odniesieniu do rozmiaru poszczególnych gospodarek. W 2009 r. wartość nominalna produktów strukturyzowanych znajdujących się w obiegu, w relacji do PKB, wyniosła zaledwie około 1,1%, co stanowi niższy poziom nie tylko w porównaniu do krajów Europy Zachodniej, ale również państw naszego regionu geograficznego⁶.

⁵ Zob. M. Dyduch, *Kształtowanie się produktów strukturyzowanych na polskim rynku finansowym*, Materiały Krakowskiej Konferencji Młodych Uczonych 2010, Sympozja i Konferencje KKMU nr 5, Fundacja studentów dla AGH, Grupa Naukowa Pro Futro, Kraków 2010.

⁶ Urząd Komisji Nadzoru Finansowego, Warszawa, październik 2010.

1.1. Lokata strukturyzowana

Lokata terminowa jest to umowa między bankiem a klientem, dotycząca lokowania środków pieniężnych, zawierana na czas określony. Bank zobowiązuje się wypłacić kapitał wraz z odsetkami na koniec okresu umowy.

Lokaty strukturyzowane często nazywane lokatami inwestycyjnymi bądź lokatami indeksowymi funkcjonują na podstawie Ustawy Prawo Bankowe. Lokata strukturyzowana stanowi depozyt bankowy z wbudowanym instrumentem pochodnym. Konstrukcja produktu jest oparta na lokacie terminowej, w której wartość wypłacanych odsetek jest uzależniona od zdarzeń występujących na rynku kapitałowym. Połączenie dwóch instrumentów finansowych (lokaty i udziału w zmianach na rynku kapitałowym) ma na celu zwiększenie opłacalności inwestowania wraz z minimalizowaniem ryzyka wynikającego z inwestycji na rynkach kapitałowych. Główną zaletą lokat strukturyzowanych jest to, że zazwyczaj banki gwarantują zwrot całości powierzzonego kapitału, dodatkową stopę zwrotu z lokaty w postaci odsetek oraz, przy sprzyjającej koniunkturze, osiągnięcie dodatkowych zysków⁷. Produkty tego typu są objęte gwarancją Bankowego Funduszu Gwarancyjnego⁸. Dochód podlega opodatkowaniu podatkiem dochodowym od osób fizycznych.

Lokaty strukturyzowane są sprzedawane w drodze subskrypcji – okres ich nabywania jest ograniczony, tzn. emitent najpierw informuje o terminie przyjmowania zapisów na lokatę strukturyzowaną, a następnie przyjmuje kapitał od inwestorów. Po zakończeniu okresu subskrypcji są ustalane końcowe parametry produktu i rozpoczyna się jego tzw. czas życia⁹. Dana lokata strukturyzowana może być oferowana w danej subskrypcji w kilku wariantach – np. inwestor może mieć możliwość wyboru lokaty z wyższym oprocentowaniem gwarantowanym przy jednocześnie mniejszym udziale w zysku, wynikającym z kształtowania się aktywa bazowego lub odwrotnie – inwestor może otrzymać lokatę z niższym oprocentowaniem gwarantowanym, ale jego udział w zysku wynikającym ze wzrostu aktywa bazowego będzie wyższy¹⁰.

W dniu rozpoczęcia trwania lokaty (po zakończeniu subskrypcji) od ulokowanego depozytu odejmowane są środki z przeznaczeniem na zakup opcji (na indeksy określone dla danej lokaty) oraz marża banku. Następnie pozostały kapitał „odpracowuje” kwotę depozytu i gwarantowane odsetki – stąd niższa w porównaniu do lokat tradycyjnych wysokość oprocentowania oraz dłuższy

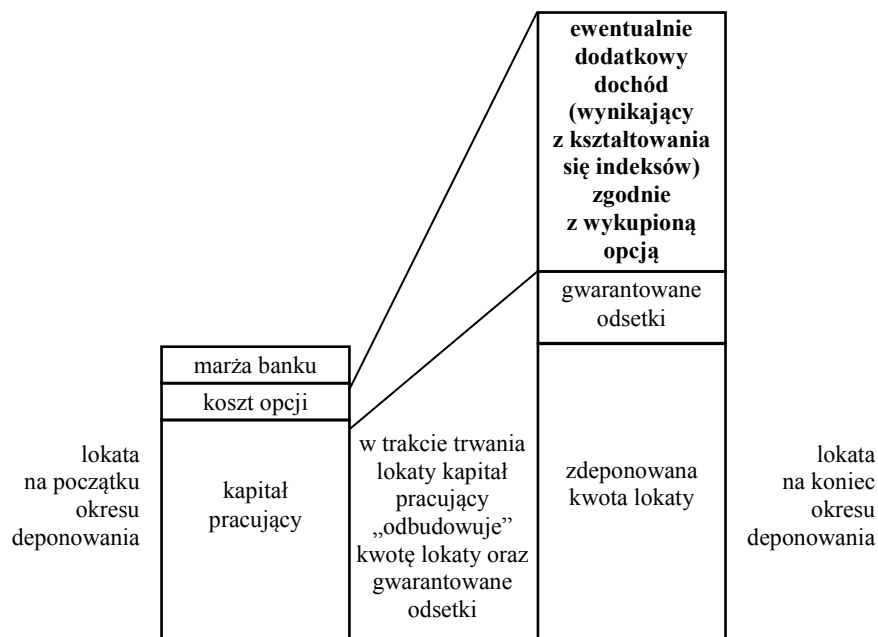
⁷ D. Korczakowski, *Lokaty inwestycyjne – źródło nieograniczonego dochodu?*, www.privatebank.pl.

⁸ Aby środki mogły być gwarantowane przez BFG, stroną umowy z klientem musi być bank, produkt musi być „emitowany” na podstawie Ustawy z dnia 29 sierpnia 1997 r. Prawo bankowe (Dz.U. 2002, nr 72, poz. 665 z późn. zm.), a „emisja” musi być czynnością bankową w rozumieniu tej ustawy.

⁹ M. Krzywda, *Produkty strukturyzowane w praktyce*, Wydawnictwo Złote Myśli & Marcin Krzywda, Gliwice 2010.

¹⁰ D. Korczakowski, op. cit.

okres trwania lokaty (kapitał musi mieć czas, by odrobić pierwotną lokatę wraz z odsetkami). O ile więc klient zerwie lokatę przed terminem, bank de facto wypłaca „odchudzony” kapitał lokaty (ewentualnie oddając część swojej marży). Produkt ten jest bardzo zyskowny dla banków – marża zwykle jest wyższa niż w wypadku tradycyjnych lokat bankowych – dodatkowo bank nie ponosi żadnego ryzyka związanego z lokatą – ewentualne ryzyko częściowo jest przerzucone na klienta. W tego typu lokatach bank nie inwestuje środków na giełdzie, a tylko wykorzystuje instrumenty pochodne z rynku kapitałowego¹¹.



Rys. 2. Schemat lokaty inwestycyjnej

Źródło: <http://www.privatebanking.pl>.

W dniu umownego zakończenia lokaty emitent-Bank wypłaca inwestorowi:

- gwarantowany zwrot kapitału,
- gwarantowane odsetki od kapitału,
- dodatkowe zyski – w wypadku ukształtowania się aktywa bazowego na odpowiednim poziomie.

¹¹ Ibid.

W przypadku gdy inwestor zerwie lokatę przed jej terminem zakończenia, to:

- otrzymuje zwrot zainwestowanego kapitału pomniejszony o ewentualną opłatę za zerwanie lokaty,
- nie otrzymuje odsetek gwarantowanych,
- nie otrzymuje dodatkowych zysków.

Lokaty strukturyzowane nie są produktem uniwersalnym i są skierowane raczej do inwestorów, którzy chcieliby zainwestować na giełdzie, ale brakuje im niezbędnej wiedzy czy doświadczenia, by inwestować samemu. Dzięki lokacie strukturyzowanej mogą oni poczuć się jak gracze giełdowi. Ponadto w odróżnieniu od bezpośrednich inwestycji giełdowych, o ile dotrzymają oni terminu umownego zakończenia lokaty, mają gwarancję zwrotu zainwestowanego kapitału. Przy korzystnej koniunkturze, zyski będą jednak raczej mniejsze niż przy samodzielnym trafnym inwestowaniu na giełdzie w warunkach hossy. Reasumując, do najważniejszych zalet lokaty strukturyzowanej można zaliczyć:

- gwarantowany zysk,
- potencjalną możliwość uzyskania znacznych zysków,
- prostotę konstrukcji – z punktu widzenia inwestora,
- inwestowanie na giełdzie bez konieczności ponoszenia opłat za zarządzanie.

Do podstawowych wad lokat strukturyzowanych można zakwalifikować:

- brak płynności,
- brak bieżącej informacji o zyskach z lokaty strukturyzowanej,
- ofertę skierowaną do ściśle określonego grona osób,
- skorzystanie z usługi jest możliwe tylko w określonym i zależnym od banku terminie,
- wysokie koszty związane z wcześniejszą likwidacją lokaty, tzn. likwidacją przed upływem terminu.

2. Cel i metodyka badania

Celem pracy jest wycena produktu strukturyzowanego. Do badania wykorzystano produkt strukturyzowany w formie lokaty strukturyzowanej opartej na kursie fixingu EUR/PLN ogłaszanym przez Narodowy Bank Polski. Zysk z analizowanej lokaty jest uzależniony od kształtowania się kursu EUR/PLN na fixingu NBP. Zysk z lokaty wyniesie min. 7,80% p.a. (ostateczna wartość zostanie ustalona w dniu 10 sierpnia 2010 r.), jeżeli kurs EUR/PLN na fixingu NBP w dniu 8 sierpnia 2011 r. będzie większy lub równy od kursu początkowego z dnia 10 sierpnia 2010 r. pomniejszonego o 0,16 zł albo będzie mniejszy lub równy od kursu początkowego powiększonego o 0,16 zł. W prze-

ciwnym wypadku zysk wyniesie 0%, a klient otrzyma zwrot zainwestowanych środków. Szczegółową charakterystykę analizowanego produktu przedstawia tabela 1.

Tabela 1

Zasady lokaty strukturyzowanej Stabilna złotówka

Nazwa lokaty	Lokata strukturyzowana Stabilna złotówka
Minimalna kwota lokaty	1000 PLN
Data rozpoczęcia okresu subskrypcji	03.08.2010
Data zakończenia okresu subskrypcji	09.08.2010
Data rozpoczęcia lokaty strukturyzowanej	10.08.2010
Data zakończenia lokaty strukturyzowanej	10.08.2011
Nazwa aktywa bazowego	Kurs fixingu EUR/PLN ogłaszany przez Narodowy Bank Polski
Gwarancja kapitału	100% ochrony kapitału na zakończenie inwestycji
Formuła obliczania odsetek od lokaty	Oprocentowanie lokaty w skali roku obliczane jest według poniższej formuły: $X \quad \text{jeżeli} \quad I_p - 0,16 \text{ PLN} \leq I_k \leq I_p + 0,16 \text{ PLN}$ 0% w przeciwnym przypadku gdzie: I_p – kurs EUR/PLN na fixingu NBP w dniu 10.08.2010 r. I_k – kurs EUR/PLN na fixingu NBP w dniu 08.08.2011 r. X – wysokość oprocentowania w skali roku. Według danych rynkowych z dnia 02.08.2009 wartość ta wyniosłaby 8,00% p.a. ostatecznie wartość zostanie ustalona w dniu 10.08.2010 i będzie nie niższa niż 7,80% p.a.
Formuła obliczania odsetek w przypadku wycofania lokaty	Wysokość odsetek w przypadku wycofania lokaty przed terminem uzależniona jest od rynkowej wyceny strategii inwestycyjnej w dniu realizacji zerwania
Oprocentowanie gwarantowane	0%
Cykl wycofania lokaty	Zerwania lokat odbywają się w cyklach tygodniowych. Cykl trwa od środy do wtorku włącznie każdego tygodnia okresu umownego
Oплата za wycofanie lokaty	Maksymalna opłata manipulacyjna to 5,90% kwoty zainwestowanych środków. Wysokość opłaty zależy od daty złożenia dyspozycji zerwania LOKATY ustrukturyzowanej zgodnie z poniższym: Dyspozycję zerwania złożono Od dnia: Do dnia: Prowizja:

Źródło: <http://www.mBank.pl>.

Z przedstawionej charakterystyki lokaty strukturyzowanej wnioskujemy, że zysk z inwestycji w analizowany produkt strukturyzowany jest zależny od kształtowania się kursu EUR/PLN. Zatem celem oszacowania ewentualnego zysku z inwestycji w lokatę strukturyzowaną Stabilna złotówka, należy właściwie oszacować tylko wartość kursu EUR/PLN na dzień 8 sierpnia 2011 r., gdyż zysk inwestora zależy od kształtowania się właśnie tego kursu.

W dalszej części artykułu skoncentrowano się na wycenie przyszłej wartości kursu EUR/PLN.

2.1. Opis badania

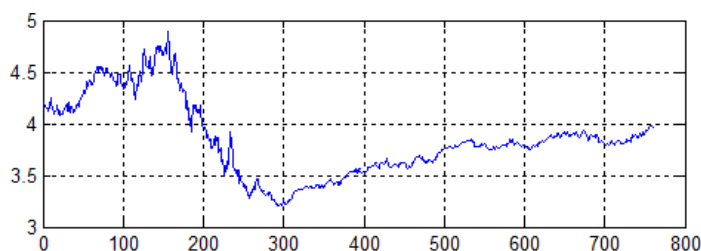
Jak wspomniano we wprowadzeniu, narzędziami zastosowanymi do oszacowania wartości kursu EUR/PLN w artykule są sieci neuronowe i falki, czyli funkcje, które dzięki przesunięciu i przeskalowaniu generują bazę ortonormalną przestrzeni $L^2(R)$. Funkcje te nazwano falkami, ponieważ są znakoziemne, czyli falują i w sposób istotny różnią się od zera tylko na małym odcinku. Falki w transformacie falkowej odpowiadają funkcjom sinus i cosinus w transformacie Fouriera, jednak te pierwsze nie są jednoznacznie zdefiniowane. Ponadto rodzina falek użyta w danej transformacie musi spełniać pewne warunki, tj. falki muszą być ortogonalne, ich średnia całkowita musi być równa 0 (falka musi „oscylować”), a w nieskończoności ich wartości dążą do zera.

Transformata falkowa:

- reprezentuje sygnał jednowymiarowy w przestrzeni i czasie,
- nadaje się przede wszystkim do analizy sygnałów niestacjonarnych,
- ma zastosowanie w analizie sygnałów z radaru, sonaru, sejsmicznych, EKG, muzycznych.

Model wykorzystany do prognozowania wartości kursu EUR/PLN można opisać następująco:

- 1) W pierwszej kolejności dokonujemy podziału szeregu przedstawionego na rysunku 3 (761 obserwacji z okresu 24.07.2007-11.08.2010 r.) na podszeregi 8-elementowe.



Rys. 3. Wartości kursu EUR/PLN wykorzystane w badaniu

- 2) Następnie dokonujemy transformaty falkowej podszeregów 8-elementowych, tzn. każdy z 500 wyodrębnionych podszeregów 8-elementowych poddajemy transformacie falkowej algorytmem „a trous”¹².

Należy w tym miejscu wspomnieć, że w odróżnieniu od transformacji Fouriera, transformacja falkowa dokonuje rozbicia sygnału na sygnały elementarne zwane falkami. Są to przebiegi ciągle oscylacyjne o różnych czasach trwania i o zróżnicowanym widmie. Aby klasa sygnałów mogła posłużyć za bazę do przekształceń falkowych, wszystkie falki muszą być wzajemnie ortogonalne. Oznacza to, że żadnej z falek nie można zapisać jako liniowej kombinacji dowolnych pozostałych ze zbioru¹³.

Wielką zaletą analizy falkowej, w stosunku do analizy częstotliwościowej, jest fakt, iż rozdzielczość czasowa transformacji może się zmieniać, ponieważ jest ona zależna od częstotliwości falki – lepsza rozdzielczość dla wyższych częstotliwości. Pełen zbiór falek użytych do dekompozycji składa się z przebiegu podstawowego oraz pozostałych przebiegów, które są jego kopiami przesuniętymi w czasie oraz rozciągniętymi lub ściśniętymi na osi czasu. Funkcja podstawowa ma zerową wartość średnią. Może nią być przykładowo przebieg oscylacyjny zmodulowany funkcją dzwonu, tj. o obwiedni symetrycznie narastającej i opadającej. Każda falka jest opisana współczynnikiem kompresji, który mówi o kompresji lub rozciągnięciu czasowym w stosunku do przebiegu podstawowego. Na podstawie wyliczonych współczynników kompresji i przesunięć czasowych wszystkich falek składających się na analizowany przebieg buduje się funkcję dwóch zmiennych opisującą czasowo-częstotliwościowe własności przebiegu. Lokalne maksima tej funkcji świadczą o pojawieniu się przejściowego przebiegu o częstotliwości odwrotnie proporcjonalnej do wartości współczynnika kompresji.

- 3) Odwrotna transformata falkowa dla kilku zbiorów 8-elementowych.

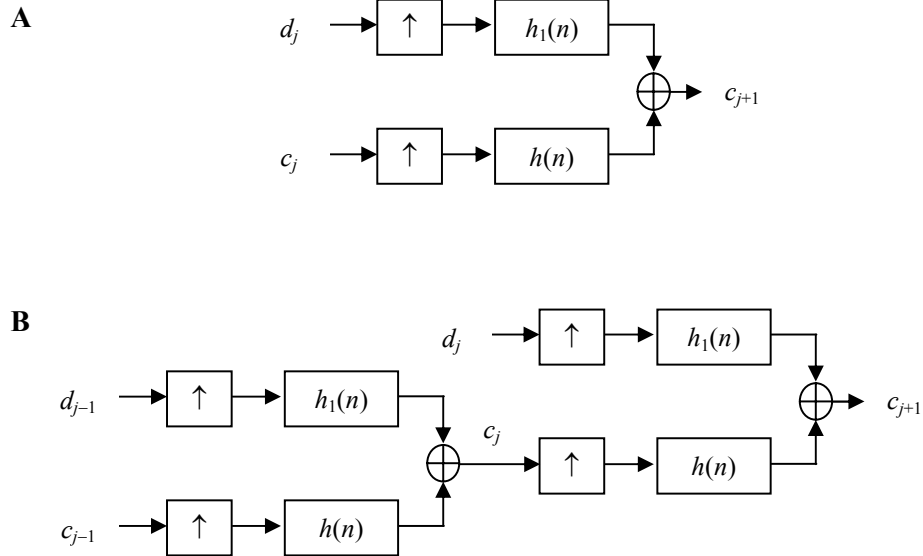
Algorytm obliczania IDWT jest następujący:

- a) określenie współczynników filtrów: dolnoprzepustowego $h(n)$ i górnoprzepustowego $h_1(n)$,
- b) uzupełnienie sygnałów aproksymacji i detali zerami,

¹² Algorytm szczegółowo opisany w pracy M. Dyduch, *Współczynniki transformaty falkowej jako narzędzie generujące prognozę przedziałową szeregów czasowych*, (w:) *Modelowanie preferencji a ryzyko '10*, (red.) T. Trzaskalik, Wydawnictwo AE, Katowice 2010.

¹³ Jaki przebieg zostanie wybrany jako falka zależy od potrzeb analizy, od tego jakie kształty rytmu są poszukiwane w zapisie. W przypadku gdy obiektem zainteresowań są przebiegi niegasnące lub trwające długo w porównaniu z oknem analizy, gdy nieistotna jest lokalizacja czasowa przebiegów przejściowych, najlepszą bazą będzie zbiór sinusoid, a więc użycie transformacji Fouriera. Jeżeli badany sygnał jest zasadniczo niestacjonarny, bogaty w przebiegi przejściowe, a przedmiotem analizy ma być lokalizacja czasowa przebiegów przejściowych o określonych częstotliwościach, bazą do analizy będą falki.

- c) dokonanie splotu uzupełnionego zerami wektora c ze współczynnikami filtra $h(n)$, otrzymując dolnoprzepustową informację o sygnale, oraz uzupełnionego zerami wektora d z $h_1(n)$, otrzymując górnoprzepustową informację o sygnale,
- d) sumowanie wektorów otrzymanych w punkcie c.

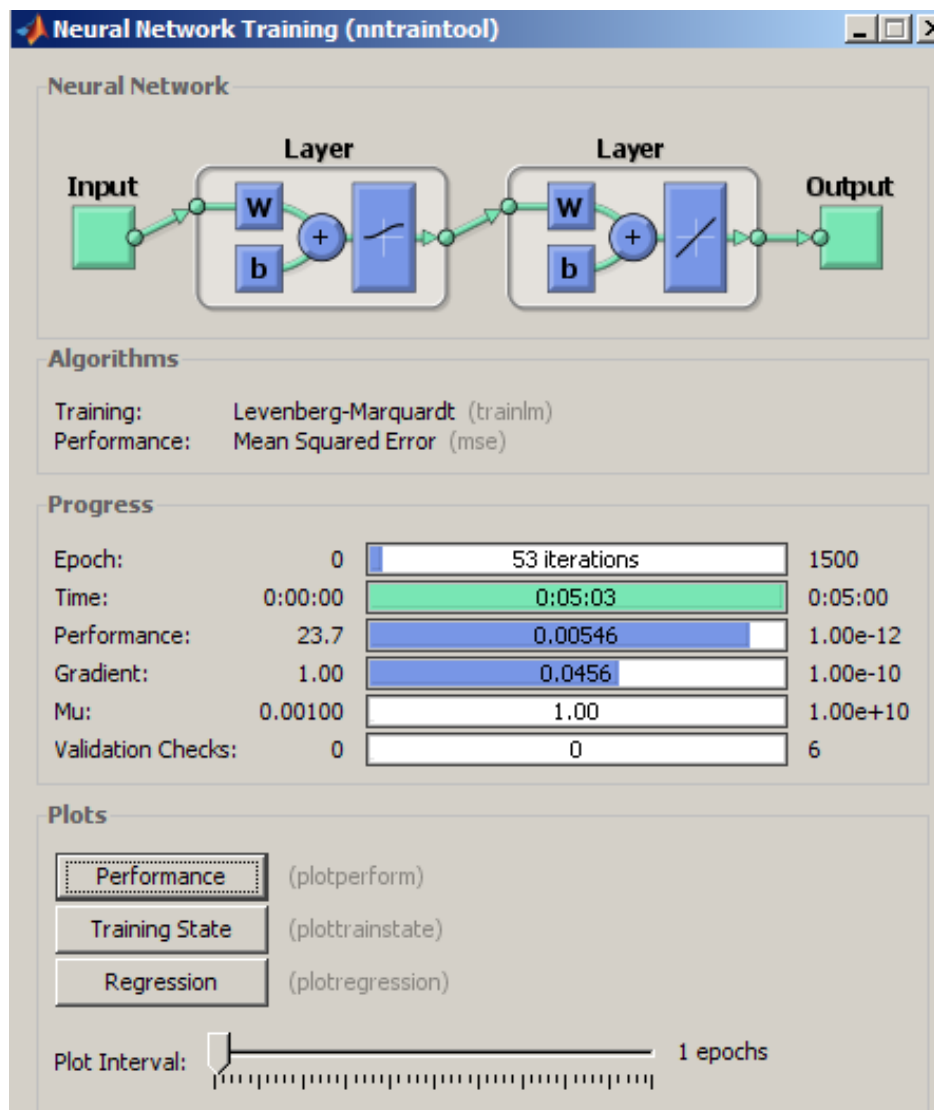


Rys. 4. A – odwrotna dyskretna transformata falkowa IDWT; B – synteza wielopoziomowa

Współczynniki filtra h_1 oblicza się ze wzoru:

$$h_1(n) = (-1)^n h(N - n).$$

- 4) Generowanie współczynników falkowych dla kolejnych chwil czasowych, czyli dla chwil prognozowanych przy użyciu sztucznej sieci neuronowej.



Rys. 5. Parametry sieci

Wprawdzie najlepiej znanym przykładem algorytmu uczenia sieci neuronowej jest metoda wstecznej propagacji błędów, jednakże nowoczesne algorytmy drugiego rzędu, takie jak: metoda gradientów sprzężonych i metoda Levenberga-Marquardta¹⁴ są w większości zastosowań istotnie szybsze (tzn. dużo szybsze – np. o rząd wielkości!) i dlatego w przeprowadzonym badaniu wykorzystano metodę *Levenberga-Marquardta*. Klasyczna metoda wstecznej propagacji błędów też ma wiele istotnych zalet, przede wszystkim jest bezspornie najprostszym algorytmem do zrozumienia. W metodzie wstecznej propagacji błędów (BP) na początku obliczany jest wektor gradientu wyżej opisanej powierzchni błędu. Wektor ten określa linię przechodzącą przez bieżący punkt i wyznaczającą kierunek, w którym spadek wartości błędu jest najszybszy. Dokładniej – gradient wyznacza kierunek najszybszego wzrostu funkcji błędu, w związku z czym kierunek zmiany wag w sieci jest dokładnie przeciwny do kierunku wyznaczanego przez gradient. Jeśli zostanie wykonany niewielki krok w wyznaczonym kierunku, to w osiągniętym w ten sposób punkcie wartość błędu będzie (zazwyczaj) mniejsza. Sekwencja takich przesunięć (o wielkości zmniejszającej się wraz z przybliżaniem się do dna „doliny” wyznaczającej minimum funkcji błędu) ostatecznie prowadzi do osiągnięcia pewnego minimum. Niestety nie zawsze jest to minimum globalne.

Istotnym problemem związanym ze stosowaniem metody BP jest określenie wielkości wykonywanych kroków. Stosowanie dużych kroków zapewnia szybszą zbieżność algorytmu, ale może również powodować przeskoczenie ponad punktem stanowiącym rozwiązanie lub też (w przypadku bardzo urozmaiconych powierzchni błędu) powodować gwałtowny krok w niewłaściwym kierunku. Klasycznym przykładem takiego zachowania się algorytmu uczenia jest sytuacja, w której obserwowany jest bardzo powolny spadek wartości błędu w trakcie uczenia, ponieważ algorytm posuwa się wzdłuż stromej, wąskiej doliny, w której nie może osiągnąć dna, tylko z uwagi na wykonywane zbyt duże kroki ciągle następują odbicia pomiędzy jej przeciwnymi zboczami. W przeciwieństwie do wyżej opisanej sytuacji, stosowanie zbyt małych kroków może wprawdzie prowadzić we właściwym kierunku, ale zbyt wolno, ponieważ wymaga stosowania bardzo dużej liczby iteracji. W praktyce rozwiązuje się zasygnalizowane problemy w taki sposób, że stosuje się wielkość kroku proporcjonalną do nachylenia (dzięki czemu algorytm zatrzymuje się w minimum) i do pewnej stałej, zwanej współczynnikiem uczenia¹⁵.

¹⁴ Ch.M. Bishop, *Neural Networks for Patterns Recognition*, Oxford University Press, Shepherd 1995.

¹⁵ Właściwe określenie wartości współczynnika uczenia jest różne w różnych aplikacjach i jest najczęściej określane na drodze eksperymentalnej. Można również stosować zmienny w czasie współczynnik uczenia, najczęściej o wartości zmniejszającej się w trakcie pracy.

- 5) Odwrotna transformata falkowa – efekt – prognozowane wartości kursu EUR/PLN; na tym etapie, wykorzystując współczynniki falkowe wygenerowane przez sieć neuronową, konstruujemy poprzez odwrotną transformatę falkową oryginalny szereg czasowy, tzn. konkretną wartość kursu EUR/PLN w preferowanym dniu.
- 6) Interpretacja otrzymanego wyniku – otrzymana prognoza jest prognozą w postaci szeregu 8-elementowego, czyli każdorazowo otrzymujemy okres 8 dni. Otrzymana prognoza dla interesującego okresu jest przedstawiona w tabeli 2.

Tabela 2

Wyniki

Data	Prognozowana wartość EUR/PLN
08.08.2011	3,9659
05.08.2011	3,9946
04.08.2011	3,9427
03.08.2011	4,0038
02.08.2011	3,9755
01.08.2011	3,981
29.07.2011	3,9654
28.07.2011	3,9755

Źródło: Opracowano na podstawie przeprowadzonych obliczeń.

3. Wnioski z przeprowadzonego badania

Z przedstawionej charakterystyki lokaty strukturyzowanej wynikało, że zysk z inwestycji w analizowany produkt strukturyzowany jest zależny od kształtowania się kursu EUR/PLN. W tym celu, aby oszacować ewentualny zysk z inwestycji w lokatę strukturyzowaną Stabilna złotówka, należało właściwie oszacować tylko wartość kursu EUR/PLN na dzień 08 sierpnia 2011 r., gdyż zysk inwestora zależy od kształtowania się właśnie tego kursu.

Taką wartość kursu EUR/PLN uzyskano na podstawie wygenerowanych przez sieć neuronową współczynników odwrotnej transformaty falkowej i wynosi ona na dzień 08 sierpnia 2010 r. 3,9659 (tabela 2).

Dysponując prawdopodobną wartością kursu EUR/PLN w dniu 8 sierpnia 2011 r., można oszacować ewentualne zyski z inwestycji w lokatę strukturyzowaną Stabilna złotówka.

Mamy zatem:

- kurs EUR/PLN z dnia 10 sierpnia 2010 r.: $3,9727 = I_p$,
- prognozowany na podstawie autorskiego modelu kurs EUR/PLN na dzień 8 sierpnia 2011 r. wynosi: $3,9659 = I_k$.

Reasumując:

$$I_p - 0,16 \text{ PLN} \leq I_k \leq I_p + 0,16 \text{ PLN}$$

$$3,9726 - 0,16 = 3,8126$$

$$3,9726 + 0,16 = 4,1326$$

Warunek zysku opisany w tabeli 1 jest zatem spełniony, więc prawdopodobnie zysk z lokaty wyniesie: $x * \text{zainwestowany kapitał}$, czyli $7,8\% * \text{zainwestowany kapitał}$.

VALUATION OF STRUCTURED PRODUCT AS AN EXAMPLE OF STRUCTURED DEPOSITS STABLE ZLOTY

Summary

Structured products are one of the most flexible modern investment instruments. This type of investment is nothing but a 'packed' in a box of investment strategy, designed based on derivatives (options and to a lesser dimension swap'y), characterized by complete or partial guarantee of the invested capital, assuming the transaction until its collapse.

Structured products offered on the Polish market may have various legal forms. Structured product may be offered in the form of bank deposits, certificate of deposit, life insurance, closed-end investment fund, a Luxembourg fund, bond or investment certificate. The article attempts to estimate the profit from investments in structured products on the Polish financial market. The object of study is structured deposit issued by mBank, whose index of fixing the basic rate of EUR/PLN announced by the NBP.

Aleksandra Szpulak*

PROGNOZOWANIE W ZARZĄDZANIU AKTYWAMI OBROTOWYMI PRZEDSIĘBIORSTWA

Wprowadzenie

Można wyróżnić dwie drogi budowy modelu prognostycznego wynikające z samej natury prowadzonych badań. Są to badania eksploracyjne (ang. *explore*) i badania wyjaśniające (ang. *explain*). W badaniach eksploracyjnych poszukuje się w szeregach czasowych oraz w związkach pomiędzy zmiennymi prawidłowości, które następnie stanowią podstawę budowy modelu prognostycznego. W badaniach wyjaśniających podstawą budowy modelu jest wiedza o samym procesie, który generuje dane. W pierwszym przypadku postępowanie rozpoczyna się od danych i prowadzi do modelu, a w drugim – od modelu do danych. W literaturze podstawy prognozowania sformułowane w wyniku badań eksploracyjnych nazywane są ontologicznymi, a w wyniku badań wyjaśniających – gnoseologicznymi¹.

W prognozowaniu finansowym w przedsiębiorstwie budowa prognoz poprzez prowadzenie badań eksploracyjnych nie przynosi zadowalających efektów. Wystarczy przeanalizować błędy prognoz analityków finansowych stosujących podejście eksploracyjne, aby się o tym przekonać. Odpowiedzialny za taki stan rzeczy jest bardzo duży składnik losowy, który utrudnia identyfikację prawidłowości stanowiących podstawę prognozowania, dlatego wyznaczone prognozy charakteryzują się niską jakością *ex post* – są nietrafne, a dodatkowo obciążone; niejednokrotnie trzeba w ogóle zrezygnować z budowy prognoz. Czy można w takiej sytuacji wskazać sposób, który umożliwi poprawę jakości budowanych prognoz?

W prowadzonych badaniach autorka postawiła hipotezę głoszącą, że wykorzystanie podstaw gnoseologicznych zamiast ontologicznych w prognozowaniu finansowym pozwoli poprawić jakość *ex post* budowanych prognoz. Choć przypuszczalnie, prognozowanie oparte na podstawach gnoseologicznych

* Uniwersytet Ekonomiczny we Wrocławiu.

¹ P. Dittmann, E. Szabela-Pasierbińska, I. Dittmann i in., *Prognozowanie w zarządzaniu przedsiębiorstwem*, Wolters Kluwer, Kraków 2009, s. 16; *Prognozowanie gospodarcze. Metody i zastosowanie*, (red.) M. Cieślak, Wydawnictwo Naukowe PWN, Warszawa 2001, s. 20.

nie pozwoli zwiększyć na tyle dopasowania modelu prognostycznego, aby prognozy finansowe okazały się trafne, ale można mieć uzasadnioną nadzieję, że będą to prognozy nieobciążone. Ta druga cecha powinna być wynikiem właściwego rozpoznania procesu generowania danych.

Weryfikacja sformułowanej hipotezy obejmuje kolejno etapy: 1) budowy modelu prognostycznego na podstawie badań wyjaśniających, 2) budowy modelu prognostycznego na podstawie badań eksploracyjnych, 3) wyznaczenie prognoz na podstawie zbudowanych modeli, 4) ocenę jakości *ex post* zbudowanych prognoz (ocena trafności i obciążoności prognoz). W artykule zostanie przedstawiony pierwszy etap – budowa modeli prognostycznych, wykorzystywanych w zarządzaniu aktywami obrotowymi przedsiębiorstwa.

1. Prognozowanie a planowanie finansowe w przedsiębiorstwie – ciąg dalszy

W artykule wygłoszonym na VIII konferencji naukowej „Prognozowanie w zarządzaniu firmą” argumentowano obszernie, że w przyjętej przez praktyków i naukowców nomenklaturze, cechami odróżniającymi prognozy finansowe od planów finansowych są stopień prawdopodobieństwa oraz stopień akceptacji skutków sądu o przyszłości. W powszechnym rozumieniu prognozami są sądy najbardziej prawdopodobne, a planami, sądy najbardziej pożądane². W myśl zasady, że prognozowanie jest narzędziem planowania, przyjęto, że sąd, który jest najbardziej prawdopodobny i jednocześnie pożądany, powinien być określany mianem prognozy. To kryterium, choć bardzo pomocne w zrozumieniu tego „co autor miał na myśli”, nie jest jednak wystarczające w prowadzeniu wywodów naukowych. Być może tym kryterium jest przedmiot prognozowania. Warto się zastanowić, co tak naprawdę można zaplanować, a co zaprognozować?

Od 5 lat dojeżdżam do pracy do Wrocławia z Opola pociągiem i to doświadczenie nauczyło mnie, a także moich współpracowników i studentów pewnego rodzaju pokory. Mogę wprawdzie zaplanować (wyłączając oczywiście zdarzenia nadzwyczajne), że punktualnie wstanę i wyjdę z domu, jednak czas mojego przyjazdu do Wrocławia jest trudny do przewidzenia. O ile planuję, że o określonej godzinie wyjdę na pociąg, o tyle godzina przyjazdu do Wrocławia jest tylko prognozą. Planuję godzinę wyjścia z domu, bo ta zależy ode mnie, a prognozuję godzinę przyjazdu do Wrocławia, bo ta ode mnie nie zależy.

² A. Szpulak, *Prognozowanie a planowanie finansowe w przedsiębiorstwie*, UE, Wrocław 2010.

Jeśli z perspektywy decydenta w procesie decyzyjnym ma się do czynienia ze zmiennymi: sterującymi, sterowanymi i niesterowanymi³, to tylko te pierwsze, jako że ich poziom jest wyznaczany mocą decyzji, są przedmiotem planowania, a pozostałe zmienne, tj. sterowana i niesterowana, są przedmiotem prognozowania. Sam fakt, że menedżer podejmuje się sterowania poziomem pewnych zmiennych, np. sprzedaży czy zysków, nie oznacza, że poziom tych zmiennych od niego zależy. Co wynika z tego, że menedżer zaplanuje, iż sprzeda 200 sztuk po 20 zł, jeśli nie znajdzie po tej cenie nabywcy? Dlatego w przypadku tych zmiennych nie powinno się mówić o planach, a jedynie o prognozach.

Przytoczone powyżej rozważania zostały zastosowane w prezentowanych poniżej koncepcjach modeli prognostycznych. Po wnikliwym przeanalizowaniu charakteru poszczególnych zmiennych przyjąłam założenie, że planowaniu podlega jedynie wielkość produkcji, bo ta faktycznie zależy od menedżera – pozostałe zmienne, takie jak zapasy materiałowe, zobowiązania i należności z tytułu dostaw i usług, wpływy i wydatki gotówkowe, są prognoząmi. Należy dodać, że podejście to jest zgodne z neoklasyczną teorią przedsiębiorstwa, w której przedsiębiorca/właściciel odpowiada na trzy podstawowe pytania: co? ile? oraz jak? produkować.

2. Modele prognostyczne wybranych składników aktywów obrotowych przedsiębiorstwa

W obszarze zarządzania aktywami obrotowymi są podejmowane decyzje dotyczące wielkości i struktury aktywów obrotowych. Na aktywa obrotowe składają się: zapasy, należności krótkoterminowe, inwestycje krótkoterminowe oraz krótkoterminowe rozliczenia międzyokresowe. Poniżej zostaną przedstawione modele prognostyczne głównych pozycji aktywów obrotowych, tj. zapasów materiałowych, należności z tytułu dostaw i usług, a także środków pieniężnych. Dodatkowo, z uwagi na fakt, że środki pieniężne są różnicą między wpływami a wydatkami gotówkowymi, zostanie przedstawiony również model zobowiązań z tytułu dostaw i usług.

³ A. Szpulak, *Proces prognozowania w przedsiębiorstwie na przykładzie prognozowania przepływów pieniężnych*, AE, Wrocław 2003.

2.1. Zapasy materiałowe

Zapasy materiałowe dzielą się na zapasy bieżące oraz zapas gwarancyjny. Zapasy bieżące powstają w wyniku różnic w intensywności strumieni dostaw (dopływu) i strumieni zużycia (odpływu) materiałów do produkcji. Zapas gwarancyjny jest wynikiem zakłóceń losowych występujących zarówno po stronie strumienia dostaw, jak i zużycia⁴.

Konstrukcja modelu prognostycznego zapasów materiałowych

Ustalenie *ex ante* poziomu zapasów materiałowych w danym dniu w przedsiębiorstwie produkcyjnym wymaga wyznaczenia: dziennego zapotrzebowania na materiały do produkcji (inaczej można mówić także o prognozie zużycia materiałów do produkcji), wielkości dostawy i zapasu gwarancyjnego oraz cyklu dostaw.

Pierwsza zmienna – dzienne zapotrzebowanie na materiały do produkcji M_t jest iloczynem wielkości produkcji Q_t , normy zużycia materiałów do produkcji m oraz ceny materiałów p_z :

$$M_t = Q_t \cdot m \cdot p_z \quad (1)$$

gdzie:

M_t – dzienne zapotrzebowanie na materiały do produkcji (zł),

Q_t – dzienna produkcja (jednostki naturalne),

m – norma zużycia materiału przypadająca na jednostkę produkcji,

p_z – jednostkowa cena produktu.

Drugą zmienną – wielkość dostawy D_k – wyznacza się na podstawie: zapotrzebowania na materiały do produkcji M_t i cyklu dostaw T^D wyrażonego w dniach. Wyznaczenie trzeciej zmiennej – tj. zapasu gwarancyjnego D_0 , odbywa się na podstawie oszacowania ryzyka wynikającego z losowych zakłóceń w intensywności strumienia dostaw i zużycia materiałów do produkcji. Wyznaczenie z kolei czwartej zmiennej – cyklu dostaw T^D łączy się z zagadnieniem określenia optymalnej wielkości dostawy, której ustalenie wymaga minimalizowania kosztów całkowitych zapasów (np. kosztów utrzymania

⁴ Podstawy nauki o przedsiębiorstwie, (red.) J. Lichtarski, AE, Wrocław 2007, s. 198.

zapasów, kosztów organizacji dostawy, kosztów zużycia). Zarówno wyznaczanie wielkości zapasu gwarancyjnego, jak i cyklu dostaw wykracza poza ramy tego artykułu. Użyteczne informacje na ten temat można znaleźć w wielu podręcznikach dotyczących zarządzania finansami przedsiębiorstwa⁵.

Założmy, że przedsiębiorstwo zamawia materiały do produkcji w stałym cyklu wynoszącym 10 dni. Dostarczone materiały są od pierwszego dnia aż do wyczerpania zapasu bieżącego zużywane w procesie produkcyjnym. Jakiej wielkości powinno być zamówienie? Przy cyklu dostaw wynoszącym 10 dni wielkość zamówienia będzie sumą zapotrzebowania na materiały do produkcji z 10 kolejnych dni. Pierwsza dostawa jest równa:

$$D_1 = M_1 + M_2 + \dots + M_{10}$$

a kolejna:

$$D_2 = M_{11} + M_{12} + \dots + M_{20}$$

Jaki będzie poziom zapasów materiałowych w przedsiębiorstwie? Można przyjąć, że poziom zapasów materiałowych w danym dniu w przedsiębiorstwie produkcyjnym jest sumą zapasu gwarancyjnego oraz zapasu bieżącego wynikającego ze zużycia materiałów do produkcji. Na koniec pierwszego dnia zapasy materiałowe będą wynosiły:

$$Z_{1,1} = D_0 + D_1 - M_1$$

a na koniec 10. dnia zapas bieżący zostanie całkowicie wyczerpany i dlatego poziom zapasów materiałowych będzie równy zapasowi gwarancyjnemu:

$$Z_{10,1} = D_0 + D_1 - M_1 - M_2 - \dots - M_{10} = D_0$$

W ogólnym przypadku, jeśli przyjmie się, że dostawy odbywają się w stałym cyklu wynoszącym T^D dni, to wielkość dostawy D_k określi wzór:

$$D_k = \sum_{t=T^D \cdot (k-1)+1}^{k \cdot T^D} M_t \quad (2)$$

gdzie:

D_k – wielkość k -tej dostawy (zł),

T^D – cykl dostaw (dni).

⁵ Zob. S. Wrzosek, *Zarządzanie finansami przedsiębiorstw*, AE, Wrocław 2006.

Posiadając już poziomy wszystkich czterech zmiennych potrzebnych do wyznaczenia prognozy poziomu zapasów materiałowych w przedsiębiorstwie produkcyjnym, tj. poziom zapotrzebowania na materiały do produkcji M_t , wielkość dostaw D_k , wielkość zapasu gwarancyjnego D_0 oraz cykl dostaw T^D , wyznacza się ich poziom na dany dzień w przedsiębiorstwie według następującej formuły:

$$Z_{t,k} = D_0 + D_k - \sum_{t=T^D(k-1)+1}^t M_t \quad (3)$$

gdzie:

$Z_{t,k}$ – poziom zapasów materiałowych z k -tej dostawy w momencie t ,

D_0 – zapas gwarancyjny.

Model zapasów materiałowych zdefiniowany równaniem (3) jest modelem prognostycznym zapasów materiałowych w przedsiębiorstwie produkcyjnym. Z tego równania wiadomo, że do wyznaczenia prognozy zapasów materiałowych w przedsiębiorstwie konieczne jest:

- 1) zaplanowanie dziennej produkcji Q_t ,
- 2) opracowanie norm zużycia materiałów do produkcji m w okresie prognozowanym,
- 3) wyznaczenie prognoz cen materiałów do produkcji p_z ,
- 4) ustalenie cyklu dostaw T^D oraz wyznaczenie wielkości dostaw D_k ,
- 5) wyznaczenie wielkości zapasu gwarancyjnego D_0 .

Przykład 1

W tabeli 1 zamieszczono planowaną dzienną produkcję przedsiębiorstwa na okres 20 kolejnych dni. Przy konstruowaniu planu produkcji założono, że produkcja wzrasta przeciętnie o tę samą wielkość. Wiadomo, że:

- norma zużycia materiałów do produkcji m wynosi 2 na jednostkę produkcji,
- prognozowana cena jednostkowa materiałów do produkcji p_z wynosi 3,
- zapas gwarancyjny D_0 wynosi 20,
- cykl dostaw T^D wynosi 10 dni.

Tabela 1

Planowana wielkość produkcji, zapotrzebowanie na materiały do produkcji, wielkość dostawy oraz prognoza zapasów materiałowych*

t	k	Q_t (j.nat.)	M_t (zł)	D_k (zł)	$Z_{t,k}$ (zł)
1	1	2	12	660	668
2	1	4	24		644
3	1	6	36		608
4	1	8	48		560
5	1	10	60		500
6	1	12	72		428
7	1	14	84		344
8	1	16	96		248
9	1	18	108		140
10	1	20	120		20
11	2	22	132	1860	1748
12	2	24	144		1604
13	2	26	156		1448
14	2	28	168		1280
15	2	30	180		1100
16	2	32	192		908
17	2	34	204		704
18	2	36	216		488
19	2	38	228		260
20	2	40	240		20

* Dane do przykładów 2 i 4.

Wyznamy kolejno zapotrzebowanie na materiały do produkcji M_t , wielkość dostawy D_k oraz prognozę zapasów materiałowych w przedsiębiorstwie $Z_{t,k}$.

W przypadku gdy $t = 10, k = 1$:

$$M_{10} = Q_{10} \cdot m \cdot p_z = 20 \cdot 2 \cdot 3 = 120$$

$$D_1 = \sum_{t=10 \cdot (1-1)+1}^{10} M_t = \sum_{t=1}^{10} M_t = 660$$

$$Z_{10,1} = D_0 + D_1 - \sum_{t=10 \cdot (1-1)+1}^{10} M_t = 20 + 660 - \sum_{t=1}^{10} M_t = 20$$

W przypadku gdy $t = 14$, $k = 2$:

$$M_{14} = Q_{14} \cdot m \cdot p_z = 28 \cdot 2 \cdot 3 = 168$$

$$D_1 = \sum_{t=10 \cdot (2-1)+1}^{2 \cdot 10} M_t = \sum_{t=11}^{20} 1860$$

$$Z_{14,2} = D_0 + D_2 - \sum_{t=10 \cdot (2-1)+1}^{14} M_t = 20 + 620 - \sum_{t=11}^{14} M_t = 1280$$

2.2. Należności z tytułu dostaw i usług

Należności w przedsiębiorstwie są wynikiem odroczenia płatności za dostarczone odbiorcom dobra i usługi⁶. Ich poziom w przedsiębiorstwie ustala się na podstawie dziennych przychodów ze sprzedaży oraz okresu inkasowania należności, który zależy od czynników zewnętrznych i wewnętrznych. Zewnętrzne czynniki to wiarygodność odbiorcy oraz czynniki otoczenia przedsiębiorstwa (np. popyt, zatory płatnicze, kryzys gospodarczy). Czynniki wewnętrzne to termin kredytu kupieckiego, stosowane konto oraz procedury windykacji.

Konstrukcja modelu prognostycznego należności z tytułu dostaw i usług

Założmy na moment, że przedsiębiorstwo prowadzi sprzedaż z 14-dniowym terminem zapłaty, a wszystkie należności są regulowane w terminie. Na koniec 13. dnia sprzedaży poziom należności w przedsiębiorstwie wynosi:

$$N_{13} = P_1 + P_2 + \dots + P_{13}$$

W kolejnym, 14. dniu zostaną spłacone należności równe przychodom ze sprzedaży z pierwszego dnia i jednocześnie powstaną nowe należności równe przychodom ze sprzedaży z 14. dnia:

$$N_{14} = P_1 + P_2 + \dots + P_{13} - P_1 + P_{14}$$

Ogólnie poziom należności będzie równy sumie dziennych przychodów ze sprzedaży z okresu:

$$N_t = \begin{cases} \sum_{t^*=t-T^N+2}^t P_{t^*} & t \geq T^N \\ \sum_{t^*=1}^t P_{t^*} & t < T^N \end{cases} \quad (4)$$

⁶ M. Sierpińska, M. Wędzki, *Zarządzanie płynnością finansową w przedsiębiorstwie*, Wydawnictwo Naukowe PWN, Warszawa 1998, s.128.

gdzie:

T^N – okres inkasowania należności,

P_{t^*} – dzienne przychody w okresie t (t oraz t^* są takimi samymi kolejnymi numerami tych samych okresów; wprowadzenie t^* do równania (4) oraz innych równań wynika z chęci zapisu równania z użyciem znaku sumy Σ).

Dzienne przychody ze sprzedaży uwzględnione w modelu (4) są natomiast wynikami planowanej dziennej produkcji oraz jednostkowej ceny sprzedaży:

$$P_t = Q_t \cdot p_s \quad (5)$$

gdzie:

Q_t – planowana dzienna produkcja (sprzedaż) wyrobów,

p_s – jednostkowa cena sprzedaży.

Model (4) jest modelem prognostycznym należności z tytułu dostaw i usług w przedsiębiorstwie. Z modelu wiadomo, że do wyznaczenia prognozy należności w przedsiębiorstwie konieczne jest wyznaczenie:

- 1) prognozy przychodów ze sprzedaży P_t ,
- 2) okresu inkasowania należności T^N w okresie prognozowanym.

Okres inkasowania należności jest na ogół odmienny od terminu określonego na fakturze. Można jednocześnie przyjąć, że termin inkasowania należności jest większy od terminu określonego na fakturze. Poprzez odroczenie terminu płatności odbiorca zaciąga tzw. kredyt kupiecki, który jest darmowym kapitałem obcym i dlatego decyzja o płatności przed terminem jest nieracjonalna (płatność przed terminem może być jednak uzasadniona, np. przy stosowaniu opustów cenowych)⁷.

Przykład 2

Przy założeniu, że cena jednostki produktu wynosi 8 i jednocześnie ilość wyprodukowana jest w całości sprzedana w tym samym okresie, wyznaczmy prognozę dziennych przychodów ze sprzedaży P_{t^*} , a następnie prognozę poziomu należności N_t . Okres inkasowania należności T^N wynosi 10 dni. Dane zawiera tabela 2.

⁷ Szerzej na temat sposobów szacowania okresu inkasowania należności zobacz w: A. Szpulak, *The Efficiency of the Receivable Rotation Ratio as a Measure of Accounts Receivable Settlement Period*, „Badania Operacyjne i Decyzje” 2010, nr 4.

Tabela 2

Planowana dzienna produkcja, prognoza dziennych przychodów ze sprzedaży
oraz prognoza należności z tytułu dostaw i usług*

t	Q_t	P_{t*}	N_t
1	2	16	16
2	4	32	48
3	6	48	96
4	8	64	160
5	10	80	240
6	12	96	336
7	14	112	448
8	16	128	576
9	18	144	720
10	20	160	864
11	22	176	1008
12	24	192	1152
13	26	208	1296
14	28	224	1440
15	30	240	1584
16	32	256	1728
17	34	272	1872
18	36	288	2016
19	38	304	2160
20	40	320	2304

* Dane do przykładu 2.

Prognoza przychodów ze sprzedaży P_t oraz prognoza poziomu należności N_t w $t = 4$ wynosi:

$$P_4 = Q_4 \cdot p_s = 8 \cdot 8 = 64$$

$$N_4 = \sum_{t^*=1}^4 P_{t^*} = P_1 + P_2 + \dots + P_4 = 160$$

Natomiast w $t = 14$:

$$P_{14} = Q_{14} \cdot p = 28 \cdot 8 = 224$$

$$N_{14} = \sum_{t^*=14-10+2}^{14} P_{t^*} = P_6 + P_7 + \dots + P_{14} = 1440$$

2.3. Zobowiązania z tytułu dostaw i usług

Zobowiązania z tytułu dostaw i usług nie są wprawdzie elementem aktywów obrotowych przedsiębiorstwa, jednak w celu prognozowania poziomu inwestycji krótkoterminowych, w tym także środków pieniężnych, istnieje konieczność wyznaczenia prognoz poziomu zobowiązań.

Zobowiązania w przedsiębiorstwie są wynikiem odroczenia terminu płatności za dostarczone do przedsiębiorstwa materiały do produkcji. Okres, na który przedsiębiorstwo zaciągnęło kredyt kupiecki, nosi nazwę okresu regulowania zobowiązań i zostanie oznaczony przez T^{zo} .

Konstrukcja modelu prognostycznego zobowiązań z tytułu dostaw i usług

Zobowiązanie, które zaciąga przedsiębiorstwo jest równe wielkości dostawy, którą dotychczas w naszych rozważaniach oznaczaliśmy przez D_k . Przy budowie modelu zobowiązań ma jednak znaczenie nie tylko wielkość zobowiązania, ale także moment powstania tego zobowiązania, dlatego istnieje konieczność przekształcenia D_k w D_t . Zgodnie z przyjętymi założeniami zobowiązanie powstaje w dniu, w którym następuje dostawa materiałów do produkcji, czyli:

$$D_t = \begin{cases} D_1 & t = 1 \\ D_k & t = (k-1)T^D + 1 \\ 0 & t \neq 1 \wedge t \neq (k-1)T^D + 1 \end{cases} \quad (6)$$

Przekształcenie D_k w D_t zgodnie ze wzorem (6) przedstawiono w tabeli 3.

Założmy na moment, że przedsiębiorstwo reguluje swoje zobowiązania w terminie 12 dni. Wielkość zobowiązania w $t = 1$ będzie równa wielkości dostawy D_1 . Cykl dostaw wynosi 10 dni, co oznacza, że po 10 dniach następuje całkowite zużycie zapasu bieżącego D_1 . Nowa dostawa jest ujmowana w księgach rachunkowych dnia następnego. W 10. dniu zobowiązania będą równe:

$$Zo_{10} = D_1 + D_2 + \dots + D_{10}$$

D_2 do D_{10} są równe 0, ponieważ zgodnie z założeniami przedsiębiorstwo nie realizuje dostaw materiałów do produkcji, czyli nie powstają żadne nowe zobowiązania, a dotychczasowe z okresu $t = 1$ nie zostało jeszcze uregulowane, dlatego poziom zobowiązań jest równy każdemu okresie od $t = 1$ do $t = 10$ D_1 . W 11. dniu w księgach jest ujmowana nowa dostawa materiałów do produkcji w wysokości D_2 , która jednocześnie jest wielkością zaciągniętego zobowiązania D_{11} . Poziom zobowiązań w 11. dniu będzie wynosił:

$$Zo_{11} = D_1 + D_2 + \dots D_{10} + D_{11}$$

Przy czym D_2 do D_{10} w dalszym ciągu są równe 0.

W 12. dniu przedsiębiorstwo reguluje zobowiązania zaciągnięte w $t = 1$ w wysokości D_1 i jednocześnie nie zaciąga nowych zobowiązań, dlatego poziom zobowiązań na koniec dnia będzie równy:

$$Zo_t = D_1 + D_2 + \dots + D_{10} - D_1 + D_{11}$$

gdzie D_2 do D_{10} oraz D_{11} są równe 0.

W ogólnym przypadku poziom zobowiązań z tytułu dostaw i usług w okresie t będzie sumą dokonanych zakupów w ciągu okresu:

$$Zo_t = \begin{cases} \sum_{t^*=t-T^{Zo}+2}^t D_{t^*} & t \geq T^{Zo} \\ \sum_{t^*=1}^t D_{t^*} & t < T^{Zo} \end{cases} \quad (7)$$

Model zdefiniowany wzorem (7) jest modelem zobowiązań z tytułu dostaw i usług, który może służyć wyznaczeniu prognozy zobowiązań. Z modelu wynika, że do wyznaczenia prognozy konieczne jest:

- 1) określenie okresu regulowania zobowiązań T^{Zo} w okresie prognozowanym,
- 2) wyznaczenie prognoz wielkości dostawy D_k oraz momentu powstania zobowiązania z tego tytułu D_t .

Przykład 3

Wyznamy prognozę zobowiązań z tytułu dostaw i usług dla danych z Przykładu 1. Założono, że okres regulowania zobowiązań wynosi 12 dni. Obliczenia znajdują się w tabeli 3.

Tabela 3

Prognoza wielkości dostawy oraz zobowiązań z tytułu dostaw i usług*

t	k	D_k (zł)	D_t (zł)	Zo_t (zł)
1	2	3	4	5
1	1	660	660	660
2	1		0	660
3	1		0	660
4	1		0	660
5	1		0	660
6	1		0	660
7	1		0	660
8	1		0	660
9	1		0	660
10	1		0	660

cd. tabeli 3

1	2	3	4	5
11	2	1860	1860	2520
12	2		0	1860
13	2		0	1860
14	2		0	1860
15	2		0	1860
16	2		0	1860
17	2		0	1860
18	2		0	1860
19	2		0	1860
20	2		0	1860

* Dane do przykładu 4.

Prognoza zobowiązań z tytułu dostaw i usług na $t = 4$ wynosi:

$$Zo_4 = \sum_{t^*=1}^4 D_{t^*} = 660$$

Natomiast na $t = 14$:

$$Zo_{14} = \sum_{t^*=14-12+2}^{14} D_{t^*} = \sum_{t^*=4}^{14} D_{t^*} = 1860$$

2.4. Środki pieniężne

Na ostatni omawiany składnik aktywów obrotowych – inwestycje krótkoterminowe – składają się środki pieniężne oraz inne krótkoterminowe aktywa finansowe i inwestycje krótkoterminowe. Na poziom środków pieniężnych w przedsiębiorstwie składają się środki utrzymywane ze względów transakcyjnych, ostrożnościowych czy spekulacyjnych. Do ustalenia optymalnego poziomu tych środków w przedsiębiorstwie z wykorzystaniem narzędzi optymalizacyjnych konieczne jest wyznaczenie kosztów całkowitych związanych z utrzymywaniem gotówki (w modelu Baumola-Allaisa-Tobina) lub dziennego zapotrzebowania na gotówkę wraz z określeniem odchyłeń od tego poziomu (w modelu Millera-Orra).

Posiadane przez przedsiębiorstwo środki pieniężne mogą być częściowo lokowane na rachunkach depozytowych lub w krótkoterminowych papierach dłużnych. Decyzje o tym, jaką część oraz na jaki okres ulokować środki pieniężne, menedżer podejmuje na podstawie krótkookresowych nadwyżek wpływów nad wydatkami.

Do zarządzania poziomem środków pieniężnych w przedsiębiorstwie wykorzystuje się m.in. budżet gotówki, który zawiera prognozy przepływów pieniężnych przedsiębiorstwa. Na prognozy przepływów pieniężnych składają się:

- prognozy wpływów,
- prognozy wydatków,
- moment wystąpienia wpływów,
- moment wystąpienia wydatków.

Konstrukcja modelu prognostycznego środków pieniężnych

Poziom środków pieniężnych jest różnicą między sumą wpływów a sumą wydatków skorygowanych o wielkość rezerwy gotówki. Poziom środków pieniężnych w okresie t wyznacza następująca formuła:

$$G_t = \sum_{t^*=1}^t CF_{t^*}^+ - \sum_{t=1}^t CF_t^- + G_0 \quad (8)$$

gdzie:

- CF_t^+ – wpływy w okresie t ,
- CF_t^- – wydatki w okresie t ,
- G_0 – rezerwa gotówki.

Osiągnięte przez przedsiębiorstwo wpływy zależą od przychodów ze sprzedaży, a ich moment wystąpienia – od okresu inkasowania należności. Wpływy są równe dziennym przychodom ze sprzedaży osiągniętym w okresie $t - T^N + 1$:

$$CF_t^+ = P_{t-T^N+1}$$

natomiast realizowane przez przedsiębiorstwo wydatki zależą od wielkości dostawy, a ich moment wystąpienia od terminu regulowania zobowiązań. Są równe realizowanym dostawom w okresie $t - T^{Zo} + 1$:

$$CF_t^- = D_{t-T^{Zo}+1}$$

Model zdefiniowany wzorem (8) jest modelem środków pieniężnych w przedsiębiorstwie, który można wykorzystać do wyznaczenia prognozy tych środków. Z modelu wynika, że do budowy prognozy środków pieniężnych konieczne jest wyznaczenie:

- 1) prognozy dziennych przychodów ze sprzedaży P_t ,
- 2) prognozy wielkości dostaw D_k oraz momentów powstawania zobowiązania D_t ,
- 3) okresu inkasowania należności T^N w okresie prognozowanym,
- 4) okresu regulowania zobowiązań T^{Zo} w okresie prognozowanym.

Przykład 4

Wyznamy prognozy wpływów i wydatków na podstawie wcześniejszych prognoz obejmujących prognozy dziennych przychodów ze sprzedaży oraz wielkości dostawy wraz z momentem powstania zobowiązania. Okres inkasowania należności wynosi 10 dni, okres regulowania zobowiązań – 12 dni, a rezerwa środków pieniężnych – 10.

Tabela 4

Prognoza dziennych przychodów ze sprzedaży oraz zobowiązań z tytułu dostaw i usług, wpływów i wydatków gotówkowych oraz poziomu gotówki*

t	D_t	P_t	CF_t^+	CF_t^-	G_t
1	660	16	0	0	10
2	0	32	0	0	10
3	0	48	0	0	10
4	0	64	0	0	10
5	0	80	0	0	10
6	0	96	0	0	10
7	0	112	0	0	10
8	0	128	0	0	10
9	0	144	0	0	10
10	0	160	16	0	26
11	1860	176	32	0	58
12	0	192	48	660	-554
13	0	208	64	0	-490
14	0	224	80	0	-410
15	0	240	96	0	-314
16	0	256	112	0	-202
17	0	272	128	0	-74
18	0	288	144	0	70
19	0	304	160	0	230
20	0	320	176	0	406

* Dane do przykładu 4.

Prognozy wpływów, wydatków oraz poziomu środków pieniężnych w $t = 12$ wynoszą:

$$CF_{12}^+ = P_{12-10+1} = P_3 = 48$$

$$CF_{12}^- = D_{12-12+1} = D_1 = 660$$

$$G_{12} = \sum_{t^*=1}^{12} CF_{t^*}^+ - \sum_{t^*=1}^{12} CF_{t^*}^- + G_0 = 16 + 32 + 48 - 660 + 10 = -554$$

Ujemny wynik oznacza, że przedsiębiorstwo, przy takich jak założone warunkach funkcjonowania w okresie $t = 12$, nie będzie miało wystarczających środków na uregulowanie zobowiązań. Taki wynik wynika głównie z faktu, że przedsiębiorstwo dysponuje zbyt małym wkładem własnym. Jednak zagadnienie bilansowania wpływów i wydatków wykracza poza ramy tego artykułu, dlatego nie będzie tu poruszane.

Podsumowanie

W artykule zaprezentowano część badań poświęconych weryfikacji hipotezy głoszącej, iż wykorzystanie podstaw gnoseologicznych zamiast ontologicznych w prognozowaniu finansowym pozwoli poprawić jakość *ex post* budowanych prognoz. Przy konstrukcji modeli prognostycznych opierano się na zgromadzonej i powszechnie dostępnej wiedzy o procesach generowania danych finansowych ukształtowanych przez zasady prowadzenia ksiąg rachunkowych, a opisanych przez nauki o zarządzaniu przedsiębiorstwem i jego finansami. Choć weryfikacja prognoz wyznaczonych na podstawie zbudowanych modeli jest kolejnym etapem badań, to już teraz można stwierdzić, że zbudowane modele mogą być użytecznymi narzędziami prognozowania finansowego na potrzeby zarządzania poziomem aktywów obrotowych w przedsiębiorstwie. Ich niewątpliwą zaletą jest to, że bazując na gnoseologicznych podstawach prognozowania, odzwierciedlają proces generowania danych finansowych.

FORECASTING FOR MANAGING OPERATIONAL ASSETS

Summary

The fundamental questions poses by managers usually concern future operational activities. The only one possible way of achieving this information is to make accurate forecasts. In general there are two parallel ways of building forecasting models – one is based on data exploration and second is based on data explanation. The hypothesis of this research assumes that it is possible to build corporate models of financial variables in respect to the nature of the relevant phenomena. To test this hypothesis Author has build models of variables included in working capital.

Elżbieta Sojka

ANALIZA PRZESTRZENNO-CZASOWA RYNKU PRACY W POLSCE ZE SZCZEGÓLNYM UWZGLĘDNIENIEM PROBLEMU BEZROBOCIA

Wprowadzenie

Zasadniczym celem artykułu jest przestrzenno-czasowa¹ analiza sytuacji na rynku pracy w Polsce (według województw). Analizie poddano zmiany w czasie podstawowych wskaźników związanych z aktywnością ekonomiczną ludności Polski, zwracając szczególną uwagę na problem bezrobocia, zwłaszcza bezrobocia trwającego ponad rok, określanego jako bezrobocie długotrwałe. Zbadano rozkłady bezrobotnych ze względu na wiek, płeć, wykształcenie i czas pozostawania bez pracy.

W drugiej części opracowania podjęto próbę analizy porównawczej województw pod względem sytuacji na rynku pracy w latach 2000 i 2008, wykorzystując w tym celu miernik syntetyczny. Pozwoliło to na liniowe uporządkowanie województw od „najlepszego” do „najgorszego” ze względu na przyjęte do badania cechy diagnostyczne.

1. Dynamika zmian aktywności ekonomicznej ludności Polski

W latach 2000-2009 nastąpiły w Polsce istotne zmiany w procesach aktywności zawodowej ludności, przejawiające się przede wszystkim w obniżeniu rozmiarów bezrobocia i wzroście liczby osób pracujących. Po okresie spadku w latach 2000-2003, liczba pracujących w Polsce zaczęła wzrastać i w 2009 r. stanowiła już 109,2% poziomu z 2000 r. W całym badanym okresie średnioroczny przyrost liczby pracujących wynosił 1%.

¹ Okres badawczy obejmuje lata 2000-2009, a wszelkie analizy były prowadzone na podstawie danych z Banku Danych Lokalnych GUS.

Tabela 1

Aktywność ekonomiczna ludności Polski
w latach 2000-2009 (dane średnioroczne)

Lata	Pracujący (w tys.)	Bezrobotni zarejestrowani (w tys.)	Aktywni zawodowo (w tys.)	W_{az} (w %)	W_z (w %)	S_b (w %)
2000	14 526	2785,0	17 311	56,6	55,0	16,1
2001	14 207	3169,0	17 376	56,3	53,5	18,2
2002	13 782	3431,0	17 213	55,4	51,7	19,9
2003	13 617	3329,0	16 946	54,7	51,4	19,6
2004	13 795	3230,0	17 025	54,7	51,9	19,0
2005	14 116	3045,0	17 161	54,9	53,0	17,7
2006	14 594	2344,0	16 938	54,0	54,5	13,8
2007	15 241	1618,0	16 859	53,7	57,0	9,6
2008	15 800	1211,0	17 011	54,2	59,2	7,1
2009	15 868	1411,0	17 279	54,9	59,3	8,2

Objaśnienia:

W_{az} – współczynnik aktywności zawodowej; W_z – wskaźnik zatrudnienia; S_b – stopa bezrobocia*.

* Zgodnie z definicjami przyjętymi przez GUS. Współczynnik aktywności zawodowej obliczono jako procentowy udział aktywnych zawodowo danej kategorii w ogólnej liczbie ludności danej kategorii. Wskaźnik zatrudnienia zdefiniowano jako procentowy udział pracujących danej kategorii w ogólnej liczbie ludności danej kategorii (wyróżnionej m.in. ze względu na płeć, wiek). Natomiast stopę bezrobocia rejestrowanego obliczano jako stosunek liczby bezrobotnych zarejestrowanych do liczby ludności aktywnej zawodowo.

Źródło: Na podstawie Banku Danych Lokalnych GUS.

W ciągu pierwszych trzech lat liczba bezrobotnych zarejestrowanych w urzędach pracy wzrastała w tempie przekraczającym średnio 300 tys. osób rocznie, osiągając w 2002 r. najwyższy poziom 3431 tys. osób. Stopa bezrobocia wyniosła wówczas 19,9%², co niestety stawiało nasz kraj w czołówce Europy pod względem rozmiarów bezrobocia. Kolejne lata zapisały się jako okres systematycznego spadku bezrobocia, na co niewątpliwie miało wpływ przystąpienie Polski do Wspólnoty Europejskiej i otwarcie tamtejszych rynków pracy.

Począwszy od 2004 r. emigracja Polaków za granicę na pobyt czasowy przybrała znaczne rozmiary. Według szacunków GUS w 2008 r. poza granicami Polski przebywało czasowo ok. 2210 tys. ludności, z czego ponad 82% emi-

² Liczona w stosunku do ludności aktywnej zawodowo (osoby pracujące i bezrobotne).

grantów przebywało na terytorium Unii Europejskiej. W 2008 r. w stosunku do początkowego okresu członkostwa Polski w UE liczba emigrantów zwiększyła się dwukrotnie³. Na podstawie analizy danych dotyczących liczby bezrobotnych oraz liczby migrujących w latach 2004-2008 można dostrzec prawidłowość polegającą na tym, że wzrostowi liczby emigrantów czasowych oraz pogłębianiu się ujemnego salda migracji stałych towarzyszył spadek liczby bezrobotnych w tym okresie⁴.

W latach 2004-2008 liczba bezrobotnych malała z roku na rok średnio o 21,7%. Największy spadek zanotowano w 2007 r. – o ponad pół miliona bezrobotnych w stosunku do roku poprzedniego. Ostatecznie, w 2009 r. zarejestrowano 1,41 mln bezrobotnych, co stanowiło nieco ponad 50% wartości z roku początkowego analizy. Symptomy światowego kryzysu gospodarczego zaczęły docierać do Polski w końcu 2008 r., co z pewnością przełożyło się na spadek globalnego popytu na pracę, a tym samym wzrost bezrobocia w 2009 r.

Obok wzrostu zatrudnienia drugim czynnikiem mającym wpływ na obniżenie poziomu bezrobocia w Polsce był spadek aktywności zawodowej ludności. Niski wskaźnik aktywności zawodowej był w znacznej mierze spowodowany niską podażą pracy osób młodych, poniżej 25. roku życia, jak też i osób z pokolenia 50+. Niska aktywność zawodowa ludzi młodych wynika z modelu edukacyjnego w naszym kraju, natomiast osób starszych – z powszechnego zjawiska wycofywania się z aktywności zawodowej. Wcześniejsze emerytury, łatwy dostęp do rent inwalidzkich skutkują tym, że przeciętny wiek wyjścia z rynku pracy jest w Polsce niższy od ustawowego wieku emerytalnego i w 2007 r. wiek ten wynosił 59,3 lata, przy średniej dla Unii Europejskiej 61,2 lata⁵. W świetle szacunków BAEL, główną przyczyną dezaktywizacji populacji 50+, poza emeryturą, jest choroba i niepełnosprawność (27% z 3 mln biernych zawodowo w wieku 55-64 lat w IV kwartale 2007 r.)⁶. Osoby w wieku 50+ są najliczniejszym beneficjentem rent z tytułu niezdolności do pracy. Jak wynika z danych ZUS, w grudniu 2008 r. ponad 52% osób w wieku 50-60 lat pobierało renty z tytułu niezdolności do pracy, a średni wiek osób pobierających renty wyniósł 56,1 lata⁷.

³ http://www.stat.gov.pl/gus/5840_3583_PLK_HTML.htm, Informacja o rozmiarach i kierunkach emigracji z Polski w latach 2004-2008, GUS, Warszawa 2009, s. 3.

⁴ W latach 2004-2008 liczba czasowych emigrantów kształtowała się odpowiednio (w tys.): 1000; 1450; 1950; 2270; 2210, natomiast saldo migracji stałych wynosiło: 2004 r. – (–9,4 tys.); 2005 r. – (–12,9 tys.); 2006 r. – (–36,1 tys.); 2007 r. – (–20,5 tys.); 2008 r. – (–15,4 tys.).

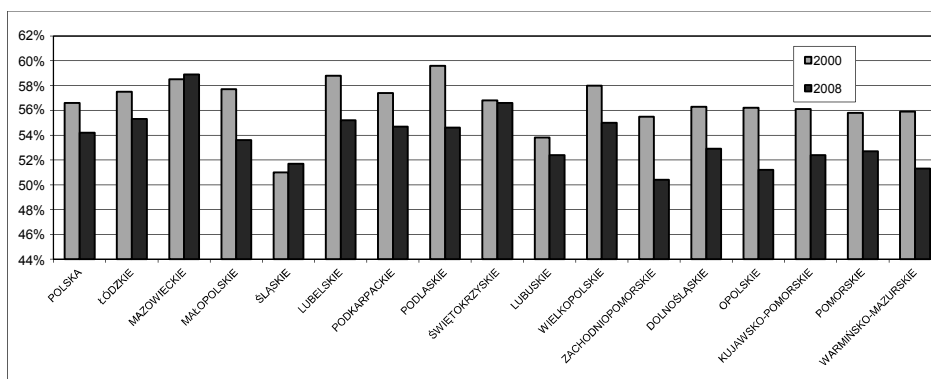
⁵ B. Kłos, *Wiek emerytalny mężczyzn i kobiet*, Biuro Analiz Sejmowych 2008, nr 3.

⁶ Aktywność ekonomiczna ludności Polski, IV kwartał 2007, GUS, Warszawa 2007, tab. 4.1, s. 161.

⁷ *Ważniejsze informacje z zakresu ubezpieczeń społecznych*, ZUS, Warszawa 2008, s. 18.

2. Aktywność zawodowa ludności (według województw)

Poziom aktywności zawodowej różni się zdecydowanie w poszczególnych województwach⁸ (rysunek 1). W 2008 r. w takich województwach, jak: mazowieckie, łódzkie, lubelskie, świętokrzyskie i wielkopolskie, poziom ten przekroczył wyraźnie poziom ogólnokrajowy, natomiast w województwach leżących przy zachodniej granicy Polski, tj. zachodniopomorskim, opolskim, śląskim oraz warmińsko-mazurskim, współczynniki aktywności zawodowej były najniższe ze wszystkich i oscylowały w granicach 50,4%-51,7%. Najwyższy współczynnik aktywności zawodowej ma województwo mazowieckie, gdzie na 100 osób w wieku produkcyjnym 59 było aktywnych zawodowo. Jest to również jedno z dwóch województw (po województwie śląskim), które odnotowało nieznaczny wzrost tego wskaźnika w stosunku do 2000 r. W pozostałych czternastu województwach mamy do czynienia ze spadkiem wartości tego wskaźnika – najwyższym pięcioprocentowym – w województwach: opolskim, zachodniopomorskim i podlaskim.

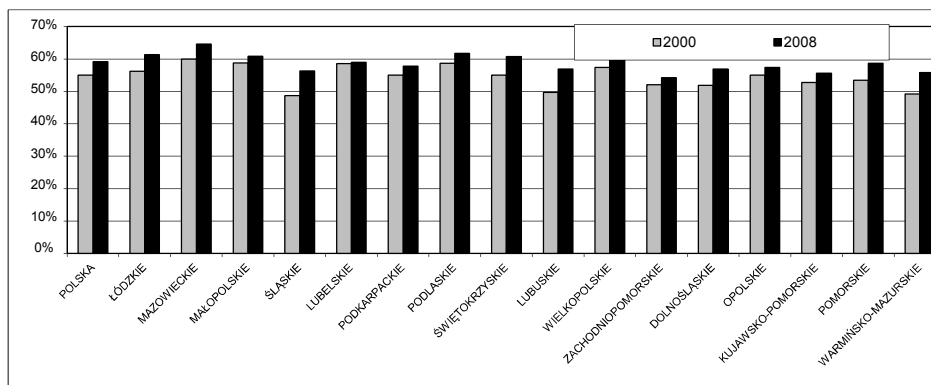


Rys. 1. Współczynniki aktywności zawodowej w województwach Polski w 2000 i 2008 r. (dane średnioroczne)

Na rysunku 2 widoczne są zmiany wskaźników zatrudnienia w wieku produkcyjnym odznaczające się wzrostem ich wartości we wszystkich województwach Polski. W zasadzie można zauważyć tę samą prawidłowość, która była widoczna w przypadku współczynników aktywności zawodowej, tzn. Zdecydowanie dominują pod względem wartości wskaźnika zatrudnienia następujące województwa: mazowieckie (64,6%), podlaskie (61,7%), łódzkie (61,3%),

⁸ Współczynnik aktywności zawodowej obliczono jako procentowy udział aktywnych zawodowo danej kategorii w ogólnej liczbie ludności danej kategorii (wyróżnionej m.in. ze względu na płeć, wiek). Wiek produkcyjny – wiek zdolności do pracy – obejmuje mężczyzn w wieku 18-64 lat i kobiety w wieku 18-59 lat.

świętokrzyskie (60,8%) i wielkopolskie (60,1%). Tam też wskaźniki przekraczają poziom odnotowany dla całego kraju. Najniższy wskaźnik zatrudnienia odnotowano w województwie zachodniopomorskim (54,2%).



Rys. 2. Wskaźnik zatrudnienia w województwach Polski w 2000 i 2008 r.

Elementem rynku pracy stanowiącym swoistą przeciwwagę aktywności zawodowej jest poziom bezrobocia. Oceny tego poziomu w ujęciu regionalnym dokonano na podstawie stóp bezrobocia ogółem oraz z uwzględnieniem płci. W dalszej części artykułu przeanalizowano natomiast struktury bezrobotnych z punktu widzenia różnych cech demograficzno-społecznych, takich jak: płeć, wiek, wykształcenie i czas pozostawania bez pracy.

Najniższą stopę bezrobocia ogółem w granicach 5,5%-6,6% miały w 2008 r. województwa: pomorskie, mazowieckie, małopolskie, podlaskie, lubuskie, wielkopolskie, śląskie i opolskie (tabela 2).

Tabela 2

Stopa bezrobocia rejestrowanego w województwach Polski
w latach 2000-2008 (dane średnioroczne)

Jednostka terytorialna	Ogółem		Mężczyźni		Kobiety	
	2000 r.	2008 r.	2000 r.	2008 r.	2000 r.	2008 r.
1	2		3		4	
POLSKA	16,1	7,1	14,4	6,4	18,1	8,0
łódzkie	16,6	6,7	15,5	6,1	17,8	7,5
mazowieckie	13,1	6,0	12,0	5,7	14,5	6,4
małopolskie	11,7	6,2	10,4	5,2	13,1	7,1
śląskie	17,5	6,6	14,6	5,8	20,9	7,5
lubelskie	14,1	8,8	14,0	8,9	14,3	8,7
podkarpackie	15,9	8,2	15,4	7,5	16,4	9,1
podlaskie	15,2	6,4	14,1	6,2	16,6	6,7

cd. tabeli 2

1	2		3		4	
świętokrzyskie	15,7	8,8	14,4	8,7	17,1	8,6
lubuskie	20,7	6,5	18,4	6,0	23,5	7,1
wielkopolskie	13,6	6,1	11,3	4,6	16,4	8,0
zachodniopomorskie	19,1	9,6	17,0	9,1	21,7	10,2
dolnośląskie	21,3	9,1	19,5	8,1	23,4	10,3
opolskie	15,4	6,6	11,9	6,0	19,7	7,3
kujawsko-pomorskie	17,8	9,1	15,6	8,3	20,3	10,1
pomorskie	16,7	5,5	13,7	4,4	20,1	6,8
warmińsko-mazurskie	23,6	7,5	21,2	6,0	26,4	9,0

Źródło: Na podstawie Banku Danych Lokalnych GUS.

Jest to nieco poniżej poziomu odnotowanego dla kraju, gdzie na 100 osób aktywnych zawodowo siedem było bez pracy lub jej poszukiwało. Widoczne są różnice w poziomie bezrobocia w ujęciu regionalnym. W 2000 r. różnica między województwem o najwyższej (warmińsko-mazurskie) i najniższej (małopolskie) stopie bezrobocia wyniosła 11,6%. Po upływie ośmiu lat różnica ta zmniejszyła się do 4,1 punktu procentowego. Największy spadek stopy bezrobocia i to bez względu na płeć osób bezrobotnych wystąpił w województwach: śląskim, lubuskim, dolnośląskim i warmińsko-mazurskim.

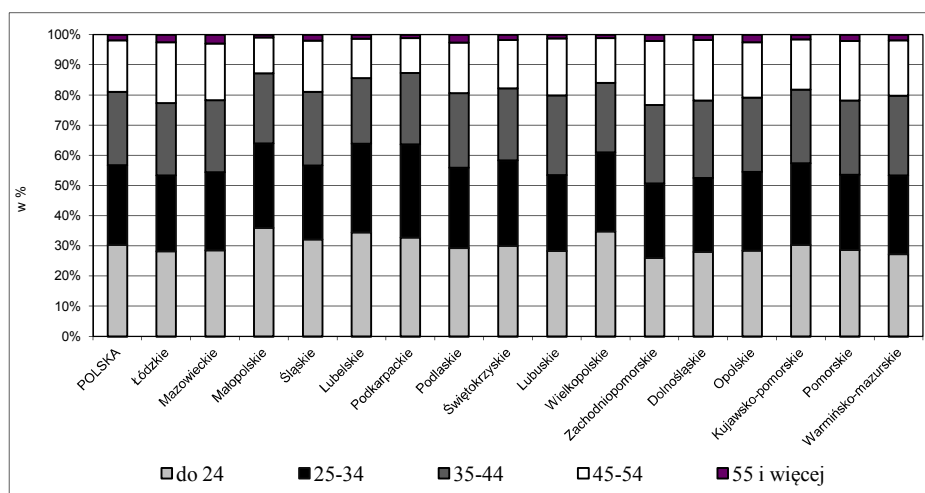
3. Struktura bezrobotnych według cech demograficznych i społeczno-ekonomicznych

We wszystkich województwach wśród bezrobotnych przeważały kobiety. Największe dysproporcje wynoszące ok. 16% w odsetkach bezrobotnych kobiet i mężczyzn w 2000 r. wystąpiły w województwach śląskim i opolskim. W 2008 r. różnica w odsetkach bezrobotnych obu płci nie uległa zmianie w województwie opolskim, natomiast zwiększyła się w województwach: śląskim, wielkopolskim i pomorskim, głównie na skutek wzrostu udziału bezrobotnych kobiet wśród wszystkich osób bez pracy. Na przestrzeni badanego okresu zagrożenie kobiet bezrobociem zwiększyło się we wszystkich województwach Polski z wyjątkiem mazowieckiego, podlaskiego i dolnośląskiego, o czym świadczy ich rosnący udział w strukturze bezrobotnych.

Powszechnie wiadomo, że sytuacja kobiet na rynku pracy jest gorsza niż sytuacja mężczyzn, a zagrożenie bezrobociem jest zdecydowanie większe. Szanse kobiet poszukujących pracy są dodatkowo zróżnicowane w zależności od wielu cech demograficznych oraz społeczno-ekonomicznych, w tym w szczególności: wieku, stanu cywilnego i wykształcenia. Kobietom jest trudniej opuścić zbiorowość bezrobotnych i częściej niż mężczyźni doświadczają bezrobocia długotrwałego.

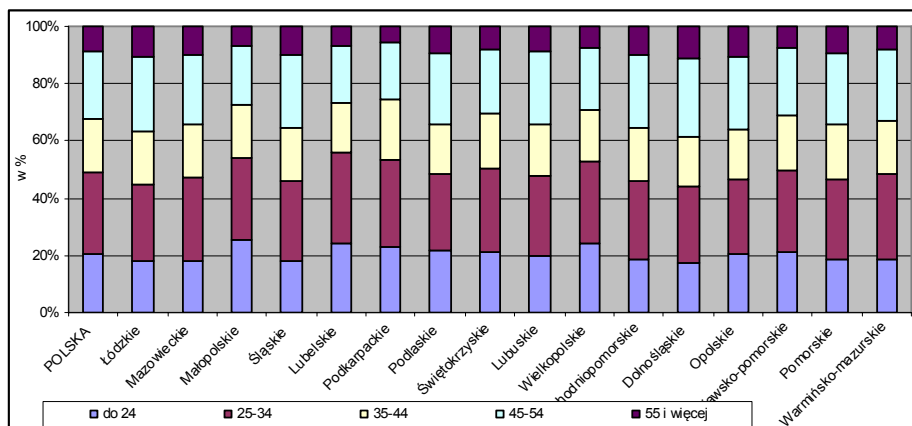
Drugą cechą charakterystyczną dla bezrobocia w Polsce, ale również w krajach europejskich, jest to, że dotyka ludzi młodych – co druga osoba bezrobotna była w wieku nieprzekraczającym 35. roku życia. Jak pisze D. Kotlorz, „(...) bezrobocie ludzi młodych jest zjawiskiem szczególnie groźnym. Przymuszona bezczynność u progu aktywności zawodowej, uniemożliwia prawidłowy rozwój społeczny i start do dorosłego życia. Brak środków do życia, niedostatek, nie spełnione aspiracje życiowe, frustracja powodują, że osoba bezrobotna jest szczególnie narażona na zjawiska patologii społecznej”⁹.

Najwyższy odsetek bezrobotnych w Polsce ogółem w 2000 r. zaobserwowano wśród młodzieży do 24. roku życia. Zasadniczy wpływ na poziom tego odsetka (30%) odegrał zwiększony napływ na rynek pracy absolwentów, tj. osób, które po zakończeniu nauki w szkole nie podjęły zatrudnienia. W 2008 r. sytuacja uległa poprawie i odsetek młodych osób bez pracy lub jej poszukujących obniżył się do poziomu 20%. Pozytywne tendencje wystąpiły również w grupie osób będących w wieku największej aktywności zawodowej, tj. 35-44 lat – spadek odsetka bezrobotnych z 24,3% w 2000 r. do 18,3% w 2008 r. Niepokojący jest natomiast wzrost tego odsetka dla osób w wieku 25-34 lat oraz 45 lat i więcej. W 2008 roku osoby bezrobotne w wieku produkcyjnym niemobilnym stanowiły prawie jedną trzecią ogółu bezrobotnych w Polsce. Podobne tendencje zmian w strukturze wieku bezrobotnych zaobserwowano w poszczególnych województwach kraju (rysunki 3 i 4).



Rys. 3. Struktura bezrobotnych w województwach Polski według wieku w 2000 r.

⁹ D. Kotlorz, U. Zagóra-Jonszta, *Rynek pracy w teorii i praktyce*, Katowice 1998, s. 108.



Rys. 4. Struktura bezrobotnych w województwach Polski według wieku w 2008 r.

Analiza powyższych wykresów pozwoliła na sformułowanie następujących spostrzeżeń.

1. Odsetek bezrobotnych w wieku do 24 lat z 2000 r. był znacznie większy niż w 2008 r., kiedy najmniejszy odsetek wystąpił w województwach: dolnośląskim (17,35%) oraz pomorskim (17,79%). Natomiast takie województwa, jak małopolskie, lubuskie i wielkopolskie, odnotowały największe udziały osób bezrobotnych w tej grupie wiekowej. To, co łączy województwa to fakt, że w miarę upływu czasu udziały bezrobotnych z najmłodszej grupy wiekowej zmniejszały się. Jest to zgodne z ogólną tendencją zaobserwowaną dla Polski.
2. W przedziale wiekowym 25-34 lat zachodziły niewielkie zmiany w czasie. Prawie w każdym z województw zanotowano nieznaczny wzrost udziałów, tylko w województwie podkarpackim frakcja bezrobotnych była w 2008 r. nieznacznie mniejsza niż w 2000 r. O ile w przypadku grupy wiekowej 18-24 lat odchylenia pomiędzy jednostkami terytorialnymi były znaczące, o tyle w przypadku bezrobotnych w wieku 25-34 lat odsetek oscylował w granicach 25%-28%. Na uwagę zasługują jedynie województwa podkarpackie oraz lubelskie, gdzie odpowiedni odsetek przekraczał 30% i był on największy spośród wszystkich w tej grupie wiekowej.
3. Nie widać większego zróżnicowania między województwami pod względem odsetka bezrobotnych w wieku 35-44 lat. Wszystkie wartości zawierają się w obszarze zmienności 17%-19% z wyjątkiem województwa pod-

karpackiego, gdzie co piąty bezrobotny był w 2008 r. w tej grupie wiekowej. Na przestrzeni lat 2000-2008 we wszystkich województwach Polski mamy do czynienia ze spadkiem udziałów – największy zaobserwowano w województwie lubuskim (z 23,8% do 18,8%), najmniejszy zaś w podkarpackim (z 23,6% do 21,3%).

4. Niepokojące jest to, że w 2008 r. bezrobocie w zdecydowanie większym stopniu niż osiem lat wcześniej, dotyczyło osób w wieku 55 lat i więcej. W zależności od województwa odpowiedni odsetek wzrósł w granicach 7%-8%. W takich województwach, jak: mazowieckie, śląskie, opolskie, dolnośląskie, zachodniopomorskie i łódzkie co dziesiąta osoba bezrobotna przekroczyła 55. rok życia. Przyczyną takiego stanu rzeczy może być brak dostatecznych kwalifikacji zawodowych osób w danym przedziale wiekowym. Ich wiedza oraz umiejętności nie są często dostatecznie dostosowane do dynamicznie rozwijającego się rynku pracy. Znaczenie mogą mieć również preferencje pracodawców, chętniej zatrudniających osoby młodsze.

W 2008 r. połowa bezrobotnych w Polsce ogółem była w wieku nie wyższym niż 35,5 lat, na co wskazuje obliczona wartość mediany. Wiek środkowy przekroczył 36 lat w województwach: łódzkim, mazowieckim, małopolskim, śląskim, lubuskim, zachodniopomorskim, dolnośląskim i opolskim. W pozostałych województwach mediana wieku oscylowała w granicach 33,9-35,8 lat. Warto zauważyć, że bezrobotni mieszkający w województwach wschodnich kraju są, średnio rzecz biorąc, młodszy niż bezrobotni zamieszkujący województwa zachodniej Polski (z wyjątkiem śląskiego), a różnica między najniższą a najwyższą wartością średniego wieku bezrobotnych w Polsce wyniosła ponad 5 lat.

Porównując struktury wieku bezrobotnych¹⁰ w poszczególnych województwach, można stwierdzić, że w 2008 r. najmniej podobne były struktury wieku bezrobotnych województwa podkarpackiego ze strukturami w województwach: dolnośląskim ($p_{ij}^* = 0,867$), pomorskim ($p_{ij}^* = 0,885$) i łódzkim ($p_{ij}^* = 0,885$). Najwyższe wartości miary podobieństwa struktur odnotowano dla następujących par województw: śląskie i zachodniopomorskie ($p_{ij}^* = 0,996$), śląskie i łódzkie ($p_{ij}^* = 0,986$) oraz śląskie i mazowieckie ($p_{ij}^* = 0,985$).

¹⁰ Podobieństwo określa się dla par struktur za pomocą miary podobieństwa: $P_{ij}^* = \sum_{k=1}^r \min(p_{ik}, p_{jk})$

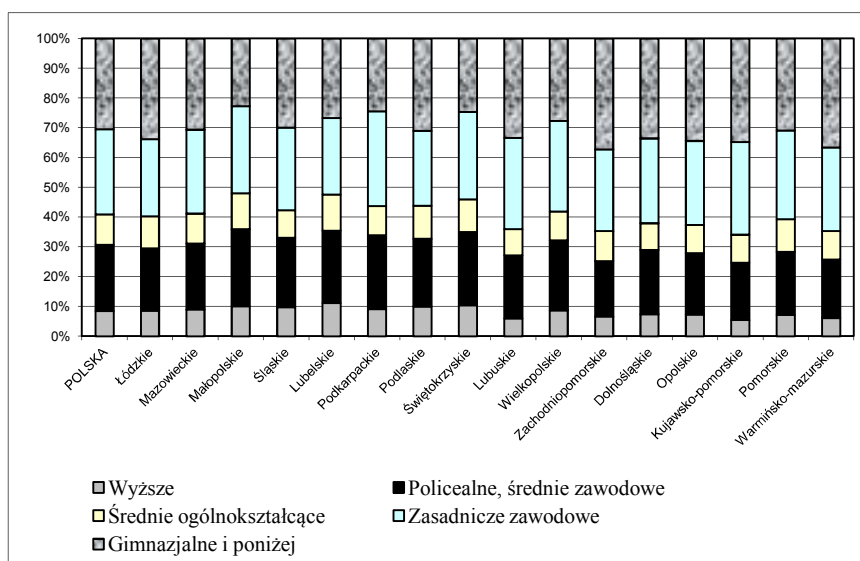
gdzie: p_{ik} – udział k -tego składnika w strukturze i -tego obiektu, p_{jk} – udział k -tego składnika w strukturze j -tego obiektu. Jeżeli struktury są całkowicie różne, to $P_{ij}^* = 0$, a jeżeli identyczne, to $P_{ij}^* = 1$.

Istotną cechą, która w znacznym stopniu decyduje o sytuacji i szansach osób bezrobotnych na rynku pracy, jest wykształcenie. Niższy poziom wykształcenia przekłada się na większe trudności ze znalezieniem zatrudnienia. Wśród osób bezrobotnych w Polsce największy odsetek stanowiły osoby z wykształceniem co najwyżej zasadniczym zawodowym, przy czym w miarę upływu czasu udział ten wykazywał tendencję spadkową z 70,4% w 2000 r. do prawie 60% w 2008 r. Spadek tego odsetka częściowo był związany z emigracją Polaków za granicę w celach zarobkowych zwłaszcza po 2004 r., kiedy to wykwalifikowani robotnicy znajdowali zatrudnienie głównie w budownictwie. Nie bez znaczenia jest również notowany od dłuższego czasu wzrost poziomu wykształcenia ludności Polski.

Zdecydowanie najniższy odsetek wśród osób bezrobotnych stanowiły osoby z wykształceniem wyższym i średnim ogólnokształcącym, które mają większe szanse na rynku pracy niż osoby bez wykształcenia lub legitymujące się jego niskim poziomem. Jest to konsekwencja zmiany świadomości bezrobotnych, którzy dążąc do zwiększenia swoich potencjalnych szans zarówno na lokalnym, jak i regionalnym rynku pracy, podejmują działania mające na celu podnoszenie swoich kwalifikacji zawodowych. Regułą jest, że kraje, w których odsetek osób z wykształceniem wyższym jest najwyższy, charakteryzują się większą wydajnością pracy, łatwiej absorbują i wdrażają nowoczesne technologie, przez co kraje te są atrakcyjnym miejscem dla nowych inwestycji. Osoby posiadające wyższe wykształcenie nawet jeżeli tracą zatrudnienie, to przeciętnie krócej będą przebywać na bezrobociu od osób słabiej wykształconych¹¹.

Niepokojący jest tu jednak dość duży wzrost w badanym okresie odsetka bezrobotnych z wykształceniem wyższym (z 2,5% w 2000 r. do 8,5% w 2008 r.). Przyczyn takiego stanu należy szukać m.in. w mało elastycznym systemie kształcenia i słabym dostosowaniu go do potrzeb zmieniającego się rynku pracy. Wśród bezrobotnych z wykształceniem wyższym największy odsetek w 2008 r. zanotowano w województwie lubelskim (11,2%) i świętokrzyskim (10,4%), natomiast najniższy w kujawsko-pomorskim (5,45%) – rysunek 5. Ze zmian, jakie zachodziły w badanych latach można wnioskować, że ukończenie studiów niestety nie gwarantuje znalezienia pracy.

¹¹ *Zatrudnienie w Polsce w 2005 r.*, (red.) M. Bukowski, MPiPS, Warszawa 2005.



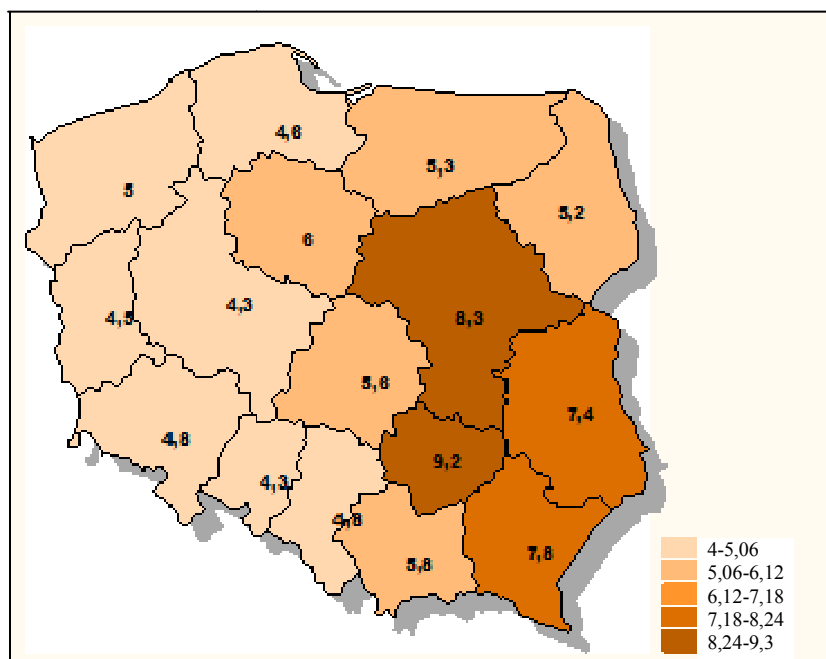
Rys. 5. Struktura bezrobotnych w województwach Polski według wykształcenia w 2008 r.

Okres poszukiwania pracy jest jednym z czynników określających sytuację na rynku pracy. Im ten okres jest dłuższy, tym mniejsze są szanse na znalezienie pracy, większe są natomiast negatywne konsekwencje dla życia osobistego osoby bezrobotnej i jego rodziny. Jeżeli czas poszukiwania pracy przekracza 12 miesięcy, bezrobocie ma charakter bezrobocia długotrwałego. W latach 2003-2008¹² struktura bezrobotnych według czasu pozostawania bez pracy w województwach Polski uległa istotnym zmianom. Zwiększył się udział bezrobotnych okresowo (tj. do 3 miesięcy), natomiast zmniejszył się odsetek bezrobotnych pozostających bez pracy powyżej 24 miesięcy. O ile w 2003 r. co piąty bezrobotny w Polsce i w przeważającej liczbie województw (z wyjątkiem łódzkiego, mazowieckiego, świętokrzyskiego i pomorskiego) pozostawał bez pracy do 3 miesięcy, o tyle pięć lat później odpowiedni odsetek oscylował w granicach od 28% w świętokrzyskim do 43% w opolskim.

Aż w siedmiu województwach: łódzkim, śląskim, wielkopolskim, zachodniopomorskim, dolnośląskim, opolskim i pomorskim odpowiedni odsetek tej grupy bezrobotnych wzrósł ponad dwukrotnie. W 2008 r. rozkład bezrobotnych według czasu pozostawania bez pracy zmienił się wyraźnie na korzyść.

¹² Analizę bezrobotnych według czasu pozostawania bez pracy przeprowadzono za lata 2003-2008 ze względu na brak danych statystycznych dla lat wcześniejszych.

Przede wszystkim zarówno w poszczególnych województwach, jak i w Polsce ogółem zmniejszył się udział osób pozostających bez pracy powyżej 12 miesięcy. Do spadku rozmiaru długotrwałego bezrobocia przyczynił się w szczególności spadek liczby osób pozostających bez pracy przez ponad 24 miesiące. Oznacza to znaczne skrócenie średniego czasu pozostawania na bezrobociu w Polsce. W 2003 r. wartość mediany tego czasu wahała się od 12 miesięcy w województwie dolnośląskim do 14,5 miesięcy w podkarpackim, natomiast w 2008 r. w niektórych województwach nastąpił ponad dwukrotny spadek średniego czasu pozostawania bez pracy, a obszar zmienności mediany mieścił się w przedziale (4,3 ; 9,2) miesięcy. Warto przy tym zauważyć, że najkrócej na bezrobociu pozostawały osoby zamieszkujące województwa leżące przy zachodniej granicy Polski, takie jak: opolskie, dolnośląskie, śląskie, lubuskie, wielkopolskie i zachodniopomorskie, gdzie średni czas pozostawania bez pracy nie przekroczył 5 miesięcy (w Polsce ogółem – 5,7) – rysunek 6.



Rys. 6. Mediana czasu pozostawania bez pracy osób bezrobotnych w województwach Polski w 2008 r.

4. Budowa syntetycznego miernika sytuacji na rynku pracy

W tej części artykułu podjęto próbę oceny przestrzennego zróżnicowania rynków pracy województw z wykorzystaniem mierników syntetycznych. Badanie przeprowadzono dla lat 2000 i 2008. Zmienne syntetyczne pozwalają na opis analizowanych jednostek, charakteryzowanych w wielowymiarowych przestrzeniach cech za pomocą jednej zmiennej agregatywnej. Umożliwia to porównanie oraz uporządkowanie tych jednostek z punktu widzenia badanego zjawiska.

Zmienne diagnostyczne będące podstawą konstrukcji miary syntetycznej powinny posiadać następujące własności:

- wysoką wartość merytoryczną,
- wysoką zdolność różnicującą analizowane jednostki terytorialne (wartość progowa współczynnika zmienności najczęściej ustalana jest na poziomie 10%),
- jednoznaczny charakter preferencji (stymulanta, destymulanta, nominanta),
- zmienne nie powinny wykazywać nadmiernego skorelowania między sobą¹³.

Ostatecznie, opierając się na kryterium merytorycznym i formalnym, w badaniu uwzględniono następujące zmienne, określając jednocześnie ich charakter:

- X_1 – wskaźnik zatrudnienia osób w wieku 55+ (stymulanta¹⁴) – zmienna ta jest miernikiem zaangażowania ludności w wieku 55+ w procesie pracy¹⁵,
- X_2 – stopa bezrobocia długookresowego w %¹⁶ (destymulanta) – zmienną charakteryzuje efektywność rynku pracy,
- X_3 – udział pracujących w usługach w % (stymulanta) – zmienna obrazuje zaawansowanie technologiczne danego województwa,

¹³ T. Grabiński, S. Wydymus, A. Zeliaś, *Metody prognozowania rozwoju społeczno-gospodarczego*, AE, Kraków 1993.

¹⁴ Wysokie wartości zmiennej świadczą o lepszej sytuacji na rynku pracy.

¹⁵ Jednym z głównych celów strategii lizbońskiej jest osiągnięcie wartości tej zmiennej w perspektywie do 2010 r. na poziomie 50%. Tylko nieliczne kraje mogą pochwalić się skuteczną polityką proaktywną i prozatrudnieniową skierowaną do najstarszych obywateli. Pozostali członkowie Wspólnoty, w tym Polska, nie zdążą z realizacją zaplanowanych celów do 2010 r. Według najnowszych założeń Ministerstwa Pracy i Polityki Społecznej z lutego 2008 r. planuje się wzrost stopy osób pracujących w wieku 55-64 lat do poziomu 40% w 2013 r., a do poziomu 50% dopiero w 2020 r.

¹⁶ Inaczej odsetek osób pozostających bez pracy 12 miesięcy i dłużej wśród ogółu aktywnych zawodowo.

X_4 – stosunek stopy bezrobocia kobiet do mężczyzn – nominanta z wartością nominalną równą jeden. Zmienna ta odzwierciedla kwestię równości płci na rynku pracy. Nadal mężczyźni mają większe szanse na rynku pracy niż kobiety, pomimo ogólnie wyższego poziomu wykształcenia wśród kobiet. Jest to jedna z najważniejszych kwestii polityki zatrudnienia również w krajach UE,

X_5 – liczba bezrobotnych na jedną ofertę pracy (destymulanta) – zmienna reprezentuje stronę popytową rynku pracy.

W kolejnym kroku analizy doprowadzono zmienne do jednorodności – wszystkie przekształcono w stymulanty. Do transformacji destymulant w stymulanty wykorzystano przekształcenie różnicowe¹⁷, natomiast przekształcając nominantę w stymulantę, zastosowano przekształcenie ilorazowe¹⁸. Celem pozabawienia miana zmiennych i ujednolicenia rzędu ich wielkości przeprowadzono dalej normalizację cech diagnostycznych, wykorzystując w tym celu przekształcenie ilorazowe:

$$z_{ij} = \frac{x_{ij}}{\max_i \{x_{ij}\}} \quad (1)$$

Wybór takiej formuły normalizacyjnej¹⁹ oznacza, że za punkt odniesienia wybrano obiekt „województwo-wzorzec” o najlepszych wartościach zmiennych, czyli taki obiekt (województwo), który jest możliwy do osiągnięcia i przyjmuje takie parametry, do których powinny dążyć inne województwa. Znormalizowane zmienne posłużyły do wyznaczenia miernika syntetycznego opisującego sytuację na rynku pracy w Polsce. Jako miarę agregującą zmienne wykorzystano średnią arytmetyczną:

$$z_i = \frac{1}{m} \sum_{j=1}^m z_{ij} \quad (2)$$

Miernik syntetyczny przyjmuje wartości z przedziału $[0,1]$. Im wyższą wartość miernika uzyskuje danych obiekt (województwo), tym lepsza w tym województwie jest sytuacja na rynku pracy.

Z przeprowadzonej analizy wynika, że istnieje duże zróżnicowanie przestrzenne kraju pod względem badanego zjawiska (tabela 3).

¹⁷ $x_{ij}^S = \max_i x_{ij} - x_{ij}^D$, gdzie x_{ij}^D – wartość j -tej destymulanty w i -tym obiekcie.

¹⁸ $x_{ij}^S = \frac{\min\{nom_j; x_{ij}^N\}}{\max\{nom_j; x_{ij}^N\}}$, gdzie x_{ij}^N – wartość j -tej nominanty w i -tym obiekcie, nom_j – nominalny poziom j -tej zmiennej.

¹⁹ Zastosowana formuła normalizacyjna powoduje unormowanie zmiennych w przedziale $[0, 1]$.

Tabela 3

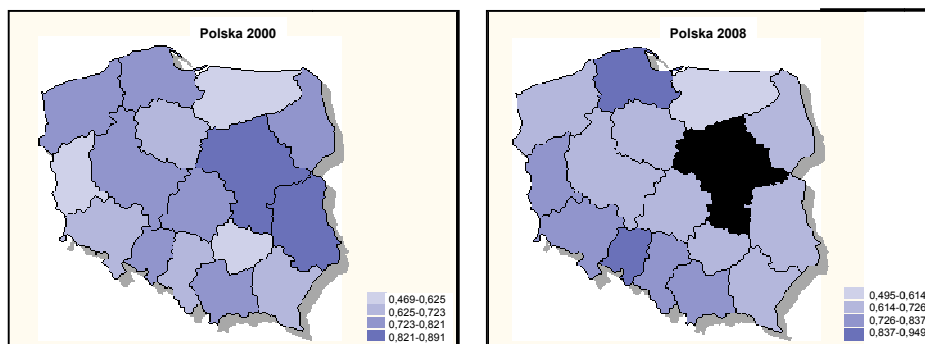
Uporządkowanie liniowe województw według wartości miernika syntetycznego w 2000 i 2008 r.

Lp.	2000 r.		2008 r.	
	Województwo	Miernik	Województwo	Miernik
1	mazowieckie	0,891	mazowieckie	0,949
2	lubelskie	0,840	pomorskie	0,880
3	opolskie	0,802	opolskie	0,839
4	małopolskie	0,790	lubuskie	0,834
5	pomorskie	0,778	śląskie	0,781
6	łódzkie	0,763	dolnośląskie	0,770
7	zachodniopomorskie	0,757	małopolskie	0,764
8	podlaskie	0,738	łódzkie	0,695
9	wielkopolskie	0,732	podkarpackie	0,692
10	śląskie	0,708	podlaskie	0,690
11	dolnośląskie	0,706	zachodniopomorskie	0,685
12	podkarpackie	0,691	lubelskie	0,670
13	kujawsko-pomorskie	0,684	wielkopolskie	0,630
14	lubuskie	0,619	kujawsko-pomorskie	0,618
15	świętokrzyskie	0,598	świętokrzyskie	0,617
16	warmińsko-mazurskie	0,469	warmińsko-mazurskie	0,495

Do województw odznaczających się stosunkowo najlepszą sytuacją na rynku pracy w 2008 r. należy zaliczyć: mazowieckie, pomorskie i opolskie. Najniższą z kolei wartość miernika syntetycznego uzyskano dla województwa warmińsko-mazurskiego. Warto przy tym zauważyć, że województwo mazowieckie wyraźnie odstaje – pod względem wartości miernika – od pomorskiego czy opolskiego i jest najbardziej zbliżone do „województwa-wzorca”.

Badanie pozwoliło również na ocenę zmian w przestrzennym zróżnicowaniu kraju pod względem analizowanego zjawiska w badanych dwóch latach. Pozytywne zmiany na rynku pracy zaobserwowano w takich województwach, jak: lubuskie (zmiana o 10 pozycji, z 14. miejsca na 4.), dolnośląskie (z 11. pozycji na 6.) i śląskie (z 10. pozycji na 5.). Zdecydowanie pogorszyła się sytuacja na rynku pracy województw: lubelskiego (spadek z 2. pozycji na 12.) oraz wielkopolskiego (z 9. pozycji na 13.). Pozostałe województwa nie zmieniły znacząco swojej pozycji. Aby ocenić stabilność i podobieństwo otrzymanych klasyfikacji, zbadano zależność w czasie za pomocą współczynnika korelacji rang Spearmana. Obliczona wartość miary korelacji równa 0,5324 świadczy o stosunkowo dużej zgodności uporządkowania województw pod względem sytuacji na rynku pracy w badanych latach.

Wykorzystując z kolei podstawowe charakterystyki opisowe miernika syntetycznego, tj. średnią arytmetyczną ($\bar{z}$) i odchylenie standardowe (S_z), dokonano klasyfikacji województw, dzieląc je na cztery grupy typologiczne: od „najlepszych”, poprzez „dobre”, „umiarkowanie dobre”, do „najgorszych” pod względem badanego zjawiska (rysunek 7).



Rys. 7. Klasyfikacja województw na podstawie syntetycznego miernika w 2000 i 2008 r.

Pierwsza grupa²⁰ obejmuje województwa: mazowieckie (z najniższą stopą bezrobocia długotrwałego – X_2), pomorskie (wysoki udział pracujących w sektorze usług – X_3) i opolskie, w którym mamy do czynienia z jedną z najniższych (po województwie śląskim) liczbą bezrobotnych przypadającą na jedną ofertę pracy – X_5 .

Cechą różnicującą czwartą grupę od pozostałych, w której znalazło się tylko województwo warmińsko-mazurskie, jest wysoka liczba bezrobotnych przypadająca na jedną ofertę pracy w tym województwie. Wartość tej zmiennej jest dziesięciokrotnie wyższa niż w województwie pomorskim i pięciokrotnie wyższa niż w mazowieckim.

Podsumowanie

W Polsce w ostatniej dekadzie wystąpiły istotne pozytywne zmiany w procesach aktywności zawodowej ludności, przejawiające się przede wszystkim w obniżeniu rozmiarów bezrobocia. Sytuacja uległa znacznej poprawie zwłaszcza po wejściu Polski do Unii Europejskiej i otwarciu tamtejszych rynków pracy. Największy spadek stopy bezrobocia i to bez względu na płeć bezrobotnych, odnotowały województwa: śląskie, lubuskie, dolnośląskie i warmińsko-mazurskie. W 2008 r. najniższą stopę bezrobocia ogółem w granicach 5,5%-6,6% miały województwa: pomorskie, mazowieckie, małopolskie, pod-

²⁰ Analiza dotyczy 2008 r.

laskie, lubuskie, wielkopolskie, śląskie i opolskie. Jest to nieco poniżej poziomu odnotowanego dla kraju, gdzie na 100 osób aktywnych zawodowo siedem było bez pracy lub jej poszukiwało.

Na podstawie przeprowadzonej analizy struktury można naświetlić sylwetkę typowego bezrobotnego w Polsce w 2008 r. Była to zwykle osoba bez stażu pracy, pozostająca na bezrobociu do 3 miesięcy, częściej była to kobieta. Jeden na trzech bezrobotnych to osoba z wykształceniem co najwyżej gimnazjalnym. Bezrobocie dotyczy głównie ludzi młodych, a połowa z nich ma nie więcej niż 35,5 lat.

W ujęciu wojewódzkim średni wiek bezrobotnego przekroczył 36 lat w województwach: łódzkim, mazowieckim, małopolskim, śląskim, lubuskim, zachodniopomorskim, dolnośląskim i opolskim. Różnica między najniższą a najwyższą wartością mediany wieku w województwach Polski wyniosła ponad 5 lat.

Badanie regionalnych rynków pracy z wykorzystaniem mierników syntetycznych pokazało, że najlepsza sytuacja na rynku pracy (z punktu widzenia przyjętych do badania cech diagnostycznych) występowała w województwie mazowieckim, a najgorsza w warmińsko-mazurskim i to zarówno w 2000, jak i 2008 r.

SPATIAL-TEMPORAL ANALYSIS OF LABOUR MARKET IN POLAND WITH SPECIAL CONSIDERATION OF THE PROBLEM OF UNEMPLOYMENT

Summary

Spatial-temporal analysis of the situation on labour market in Poland in province approach is the main purpose of the paper. Time changes in fundamental rates related to economic activity of the population of Poland with special consideration of the problem of unemployment, particularly the long-term unemployment, were analysed. Distribution of unemployed by age, gender, education and period of unemployment were studied. On the grounds of performed analysis, a profile of a typical unemployed person in Poland can be presented. Usually, it is a person without job seniority, more often a woman. One out of three unemployed people is a person with junior high school education at the most. Unemployment mainly refers to young people, and in 2008 a half of them were no more than 35.5 years old.

In the second part of the article, an attempt to perform a comparative analysis of provinces with respect to situation on labour market in 2000 and 2008 was made with application of a synthetic indicator for this purpose. This allowed for linear arrangement of provinces from „the best” to „the worst”, according to diagnostic features selected for the study. The study showed that the best situation on the labour market was observed in Mazowieckie province, and the worst in Warmińsko-Mazurskie province, in both studied years.

Jan Acedański

RYNEK AKCJI A KOSZTY WAHAŃ KONIUNKTURALNYCH W POLSCE

Wprowadzenie

W pracy z 1987 r. R. Lucas zdefiniował koszt wahań koniunkturalnych jako procentowe zwiększenie konsumpcji, które jest konieczne, aby użyteczność strumienia bieżącej konsumpcji była równa użyteczności konsumpcji w hipotetycznej gospodarce pozbawionej wahań koniunkturalnych¹. Jednym z podstawowych problemów związanych z tą definicją jest sposób pomiaru użyteczności. Lucas założył, że reprezentatywny konsument cechuje się funkcją oczekiwanej użyteczności o stałym współczynniku względnej awersji do ryzyka i oszacował, że koszt wahań nie przekracza 0,1% rocznego poziomu konsumpcji. Inni badacze, stosując funkcje użyteczności Epsteina–Zina albo wprowadzając dodatkowo do funkcji użyteczności przyzwyczajenia, uzyskali znacznie wyższe wyniki – dochodzące nawet do kilkunastu procent².

Niedawno Alvarez i Jermann zaproponowali nieparametryczną metodę pomiaru kosztów wahań strumienia konsumpcji³. Nie wymaga ona dokładnej specyfikacji funkcji użyteczności, ale o sposobie wartościowania konsumpcji wnioskuje pośrednio, badając wahania cen papierów wartościowych. W tym celu autorzy zdefiniowali krańcowy koszt wahań, który wskazuje zysk osiągany dzięki niewielkiemu zmniejszeniu wahań. Można go wyznaczyć jako iloraz cen dwóch hipotetycznych papierów wartościowych: jednego, którego wypłaty są utożsamiane ze strumieniem konsumpcji pozbawionej wahań, oraz drugiego, którego stopy zwrotu są równe stopom wzrostu konsumpcji w badanym okresie czasu. W ten sposób problem kosztów wahań sprowadza się do problemu z zakresu wyceny aktywów.

¹ R. Lucas (jr), *Models of Business Cycles*, Basil Blackwell, New York 1987.

² J. Acedański, *Metody szacowania kosztów wahań koniunkturalnych*, Studia Ekonomiczne, nr 46, AE, Katowice 2007, s. 9-27.

³ F. Alvarez, U. Jermann, *Using Asset Prices to Measure the Cost of Business Cycles*, „Journal of Political Economy” 2004, Vol. 112, No. 6, s. 1223-1256.

Do wyceny takich hipotetycznych aktywów, które nie są przedmiotem obrotu na rynku, zastosowano metodę zaproponowaną przez Cochrane'a i Saa-Requejo opartą na założeniu braku arbitrażu⁴. Ponieważ nie jest możliwe dokładne określenie ceny aktywa, które nie może być dokładnie replikowane przy pomocy istniejących na rynku aktywów bazowych, podejście to pozwala na wyznaczenie ceny średniej oraz minimalnej takiego papieru wartościowego, nakładając jedynie ograniczenie na wielkość wskaźnika Sharpe'a.

Należy podkreślić również, że w omawianej metodzie rozróżnia się koszty wszystkich wahań strumienia konsumpcji, od kosztów wahań związanych z wahaniami związanymi z cyklem koniunkturalnym, które w pracy, dla gospodarki polskiej, utożsamia się z fluktuacjami o okresie wahań nie większymi niż 24 kwartały.

W pracy zastosowano metodę Alvaraza–Jermanna do badania kosztów wahań koniunkturalnych w Polsce. Zagadnienie analizowano z dwóch względów. Po pierwsze, aby ocenić przy pomocy nowej metody skalę kosztów wahań konsumpcji związanych z cykliczną zmiennością gospodarki w Polsce. Dotychczas w literaturze krajowej zagadnienie ilościowej oceny kosztów wahań nie było podejmowane. Tymczasem jest ono fundamentalne z punktu widzenia projektowania założeń polityki makroekonomicznej i odpowiedzi na pytanie, czy powinna się skupiać na eliminacji wahań koniunkturalnych – jeśli są z nimi związane wysokie koszty społeczne – czy też powinna raczej w pierwszej kolejności sprzyjać wysokiej stopie długookresowego wzrostu gospodarczego. Po drugie, praca ma również na celu sprawdzenie samej metody. Jej autorzy zaznaczają bowiem, że uzyskana przez nich formuła kosztów wahań jest dość wrażliwa na sposób pomiaru wchodzących w jej skład zmiennych i może dawać różne rezultaty w zależności od wykorzystanych danych. Dlatego też w swojej pracy wykorzystali oni dość zróżnicowane zbiory danych. W tej pracy testuje się, jakie wyniki zostaną uzyskane dla danych pochodzących z polskiej gospodarki i polskiego rynku finansowego, które wyraźnie różnią się od danych amerykańskich. W przypadku bowiem gdyby oszacowane dla Polski koszty różniły się od wyników dla USA zdecydowanie, np. o kilka rzędów wielkości, wtedy uzasadnione byłoby podejrzenie, że jest to efekt wrażliwości samej metody.

Praca składa się z czterech części. W pierwszej przedstawione są koncepcje całkowitego i krańcowego kosztu wahań. Następnie omówiono metodologię wyceny aktywów hipotetycznych stosowanych do obliczania krańcowego kosztu wahań. W części trzeciej dyskutowany jest problem pomiaru kosztu wahań koniunkturalnych wyodrębnianych przy pomocy filtrów pasmowo-przepustowych. Wyniki badań dla gospodarki Polski zawarte są w części czwartej.

⁴ J. Cochrane, J. Saa-Requejo, *Beyond Arbitrage: Good-Deal Asset Price Bounds in Incomplete Markets*, „Journal of Political Economy” 2000, Vol. 108, No. 1, s. 79-119.

1. Krańcowy koszt wahań

Alvarez i Jermann definiują całkowity koszt wahań $\Omega(\alpha)$ następująco:

$$U((1 + \Omega(\alpha))\{C\}) = U((1 - \alpha)\{C\} + \alpha\{\bar{C}\}) \quad (1)$$

gdzie U oznacza funkcję użyteczności, $\{C\} \equiv \{C_t\}_{t=0}^{\infty}$ jest strumieniem bieżącej konsumpcji, $\{\bar{C}\} \equiv \{\bar{C}_t\}_{t=0}^{\infty}$ oznacza strumień konsumpcji pozbawionej wahań, natomiast parametr α decyduje, jaka część konsumpcji aktualnej została zastąpiona konsumpcją wygładzoną. W efekcie $\Omega(\alpha)$ pokazuje procentowy wzrost konsumpcji niewygładzonej, który w kategoriach użyteczności równoważyłby wygładzenie $100\alpha\%$ wahań konsumpcji. Całkowity koszt wahań $\Omega(1)$, rozpatrywany przez Lucasa, w tym przypadku jest równy:

$$U((1 + \Omega(1))\{C\}) = U(\{\bar{C}\}) \quad (2)$$

gdzie $\{\bar{C}\} = E(\{C\})$, przy czym E jest operatorem wartości oczekiwanej.

Zakładając, że U jest addytywną funkcją użyteczności chwilowej^{*} u , czyli:

$$U(\{C\}) = E_0 \left[\sum_{t=0}^{\infty} \beta^t u(C_t) \right] \quad (3)$$

definicję (1) można zapisać jako:

$$E_0 \left[\sum_{t=0}^{\infty} \beta^t u((1 + \Omega(\alpha))C_t) \right] = E_0 \left[\sum_{t=0}^{\infty} \beta^t u((1 - \alpha)C_t + \alpha\bar{C}_t) \right] \quad (4)$$

Obliczając pochodną tego wyrażenia ze względu na α w punkcie $\alpha_0 = 0$ oraz korzystając z faktu, że $\Omega(0) = 0$, otrzymujemy:

$$\Omega'(0) = \frac{E_0 \left[\sum_{t=0}^{\infty} \beta^t u'(C_t) \bar{C}_t \right]}{E_0 \left[\sum_{t=0}^{\infty} \beta^t u'(C_t) C_t \right]} - 1 \quad (5)$$

^{*} Założenie to nie jest konieczne. Poniższe przekształcenia są również prawdziwe bez wprowadzania założeń dotyczących postaci funkcji użyteczności.

Wyrażenie $E_0 \left[\sum_{t=0}^{\infty} \beta^t \frac{u'(C_t)}{u'(C_0)} x_t \right]$ to cena w okresie 0 aktywa, które w każdym okresie wypłaca x_t jednostek konsumpcji⁵. Stąd:

$$\Omega'(0) = \frac{P_0(\{\bar{C}\})}{P_0(\{C\})} - 1 \quad (6)$$

gdzie $P_0(\{\bar{C}\})$ to cena papieru wartościowego uprawniającego do strumienia konsumpcji pozbawionej wahań, a $P_0(\{C\})$ to cena papieru wartościowego uprawniającego do strumienia konsumpcji bieżącej.

Alvarez i Jermann pokazali, że krańcowy koszt wahań $\Omega'(0)$ jest górnym oszacowaniem całkowitego kosztu wahań $\Omega(1)$ ⁶. Na przykład dla addytywnej funkcji oczekiwanej użyteczności koszt całkowity jest równy połowie kosztu krańcowego, czyli $\Omega(1) \cong 0,5\Omega'(0)$.

2. Wyznaczanie cen aktywów hipotetycznych

W celu wyznaczenia cen $P_0(\{\bar{C}\})$ oraz $P_0(\{C\})$ zakłada się, że konsumpcja w kolejnych okresach charakteryzuje się stałą stopą wzrostu g z losowymi, nieskorelowanymi odchyleniami, których rozkład nie zależy od czasu:

$$\frac{C_{t+1}}{C_t} = \tilde{y} \quad (7)$$

przy czym $E(\tilde{y}) = 1 + g$. Ponadto konsumpcja wygładzona jest ustalona na poziomie wartości oczekiwanej konsumpcji rzeczywistej, czyli:

$$\bar{C}_t = E[C_t] \quad (8)$$

a stopa wolna od ryzyka jest stała i równa r_f .

Przy tych założeniach poszukiwane ceny aktywów można przedstawić jako:

$$P_0(\{\bar{C}\}) = \sum_{t=1}^{\infty} \frac{\bar{C}_t}{(1+r_f)^t} = \sum_{t=1}^{\infty} \frac{C_0(1+g)^t}{(1+r_f)^t} = C_0 \frac{1+g}{r_f - g} \quad (9)$$

⁵ Por. np. J. Cochrane, *Asset Pricing. Revised Edition*, Princeton University Press, Princeton 2005, s. 4.

⁶ F. Alvarez, U. Jermann, op. cit., s. 1247.

$$P_0(\{C\}) = \sum_{t=1}^{\infty} E \left[\frac{C_t}{(1+r)^t} \right] = \sum_{t=1}^{\infty} \frac{C_0(1+g)^t}{(1+r)^t} = C_0 \frac{1+g}{r-g} \quad (10)$$

gdzie r_f oraz r oznaczają stopy zwrotu w terminie do wykupu obu papierów wartościowych*. Koszt wahań równy jest więc:

$$w_0 = \Omega'(0) = \frac{r-g}{r_f-g} - 1 \quad (11)$$

Główna trudność polega na wyznaczeniu stopy zwrotu w terminie do wykupu r papieru wartościowego uprawniającego do bieżącej konsumpcji. Taki papier nie jest bowiem przedmiotem obrotu na rynku. Korzystając z założenia, że stopy wzrostu konsumpcji bieżącej mają w każdym okresie identyczny rozkład, wystarczy wyznaczyć cenę q papieru wartościowego, który w przyszłym okresie generuje wypłatę równą $\tilde{y}$. Zyskowność aktywa, które w każdym okresie uprawnia do wypłaty $\tilde{y}$, można obliczyć korzystając ze wzoru:

$$r = \frac{1+g+q}{q} \quad (12)$$

W najprostszym przypadku cena takiego hipotetycznego aktywa równa jest cenie takiego portfela aktywów bazowych, którego wypłaty w największym stopniu zbliżone są do wypłat generowanych przez wyceniany papier wartościowy. Jeżeli przez $\mathbf{x}$ oznaczymy wektor stóp zwrotu I aktywów bazowych, wśród których jest również stopa wolna od ryzyka, to poszukujemy takiego portfela $\mathbf{b}^T \mathbf{x}$, który generuje wypłaty zbliżone do realizacji zmiennej losowej $\tilde{y}$. Ponieważ ceny tych aktywów są znormalizowane do 1, więc cena takiego portfela równa jest $q = \mathbf{b}^T \mathbf{1}$. Mając n obserwacji stóp zwrotu $\mathbf{x}$ oraz y , zestawionych odpowiednio w postaci wektora $\mathbf{Y}$ oraz macierzy $\mathbf{X}$, q można szacować jako:

$$q = \mathbf{1}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \quad (13)$$

Uzyskana w ten sposób cena jest tylko przybliżeniem, ponieważ oszacowany portfel jedynie częściowo replikuje wypłaty generowane przez wyceniany papier. Cochrane i Saa-Requejo podali sposób wyznaczania górnych i dolnych ograniczeń cen takich aktywów. Szczególnie przydatne jest wyznaczenie ceny

* Ponieważ papier wartościowy uprawniający do konsumpcji wygładzonej jest pozbawiony ryzyka, dyskontowany musi być przy stopie wolnej od ryzyka.

minimalnej $\underline{q}$. Wtedy bowiem stopa zwrotu w terminie do wykupu wyznaczona zgodnie ze wzorem (12) jest maksymalna, a obliczony na jej podstawie koszt wahań (11) jest górnym oszacowaniem kosztu fluktuacji.

Zaproponowane przez nich podejście polega na znalezieniu minimalnej wartości czynnika dyskontującego m przy ograniczeniu nałożonym na jego zmienność, a dokładniej na wartość odpowiadającego mu wskaźnika Sharpe'a. Formalnie cena minimalna papieru wartościowego reprezentującego aktualną konsumpcję jest rozwiązaniem problemu:

$$\underline{q} = \min_m E[m\tilde{y}] \quad (14)$$

przy danych ograniczeniach: $\mathbf{p} = E[m\mathbf{x}]$, $m \geq 0$ oraz $D(m)/E(m) \leq h$. Pierwsze ograniczenie jest definicją czynnika dyskontującego, który jest zmienną losową wskazującą, w jaki sposób konsument wycenia wypłaty generowane przez dany papier wartościowy; $\mathbf{x}$ jest wektorem wypłat, natomiast $\mathbf{p}$ – wektorem odpowiadających im cen aktywa. Zgodnie z teorią CAPM czynnik dyskontujący równy jest stopie wzrostu użyteczności krańcowej. Ograniczenie drugie wskazuje, że użyteczność krańcowa konsumpcji jest nieujemna. Ograniczenie trzecie dotyczące zmienności czynnika dyskontującego implikuje ograniczenie wskaźnika Sharpe'a dla rynkowej ceny ryzyka, czyli stopy zwrotu z wycenianego aktywa pomniejszonej o stopę wolną od ryzyka.

Oszacowanie na podstawie próby rozwiązania problemu (14) jest postaci:

$$\underline{q} = q - \sqrt{A^2 - \mathbf{1}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{1}} \sqrt{\mathbf{Y}^T \mathbf{Y} - \mathbf{Y}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}} \quad (15)$$

gdzie $A^2 = (1 + h^2)(n(1 + r_f)^2)^{-1}$. Alvarez i Jermann pokazali, że cena minimalna tym bardziej odbiega od ceny q , im wyższy jest dopuszczalny wskaźnik Sharpe'a oraz im gorsze jest dopasowanie portfela aktywów bazowych do wycenianego papieru wartościowego.

Ceny q oraz $\underline{q}$ można również szacować przy uchyleniu założenia o jednakowym rozkładzie stóp wzrostu konsumpcji w kolejnych okresach czasu, zakładając, że konsumpcja jest modelowana przy pomocy przełącznikowego k -stanowego modelu Markowa o macierzy przejścia $\mathbf{P}$. Wtedy ceny w i -tym stanie są równe:

$$q_i = \frac{v_i}{1 + v_i} \quad \text{oraz} \quad \underline{q}_i = \frac{\underline{v}_i}{1 + \underline{v}_i} \quad (16)$$

gdzie v_i jest rozwiązaniem układu równań funkcyjnych będących odpowiednikami równania (13):

$$v_i = \sum_{j=1}^k n_j (\mathbf{X}_j^T \mathbf{X}_j)^{-1} p_{ji} \cdot \sum_{j=1}^k n_j^{-1} \mathbf{X}_j^T \mathbf{Y}_j v_j p_{ji}, \quad i = 1, 2, \dots, k \quad (17)$$

natomiast $\underline{v}_i$ jest rozwiązaniem problemu analogicznego do (14):

$$\underline{v}_i = \min_{\mathbf{m}} \sum_{j=1}^k n_j^{-1} \mathbf{m}_j^T \mathbf{Y}_j \underline{v}_j p_{ji}, \quad i = 1, 2, \dots, k \quad (18)$$

przy warunkach $\mathbf{1} = \sum_{j=1}^k n_j^{-1} \mathbf{m}_j^T \mathbf{X}_j p_{ij}$ oraz $\sum_{j=1}^k n_j^{-1} \mathbf{m}_j^T \mathbf{m}_j p_{ij} \leq (h^2 + 1)(1 + r_f)^{-2}$.

$\mathbf{X}_j$ oraz $\mathbf{Y}_j$ są macierzami obserwacji należących do j -tego stanu. Mają one wymiary odpowiednio $n_j \times I$ oraz $n_j \times 1$; $\mathbf{m}_j$ jest wektorem realizacji czynnika dyskontującego w kolejnych okresach stanu j ; p_{ji} są odpowiednimi prawdopodobieństwami z macierzy $\mathbf{P}$. Ceny q oraz $\underline{q}$ można obliczyć jako średnie z cen q_i oraz $\underline{q}_i$ ważone częstościami pojawiania się danych stanów.

3. Wyznaczanie kosztów wahań koniunkturalnych

Przedstawione powyżej podejście dotyczyło szacowania kosztów odchylenia konsumpcji od trendu, który został zgodnie z założeniem (8) ustalony na poziomie wartości oczekiwanej konsumpcji i wzrastał ze stałą stopą g . W ten sposób zostały oszacowane koszty wszystkich wahań. Jednak w rzeczywistości sam trend również podlega powolnym fluktuacjom i dlatego nie wszystkie odchylenia od ścieżki charakteryzującej się stałą stopą wzrostu można utożsamiać z wahaniami koniunkturalnymi. Alvarez i Jermann pokazali, że w przypadku gdy trend zostanie zdefiniowany jako średnia ważona obecnej i przeszłej konsumpcji:

$$\bar{C}_t = \sum_{i=0}^s a_i (1+g)^i C_{t-i}, \quad \sum_{i=0}^s a_i = 1 \quad (19)$$

to koszt odchylenia konsumpcji od tego trendu utożsamiany z kosztem wahań koniunkturalnych wyraża się wzorem*:

$$w_k = \sum_{t=1}^{\infty} \frac{r-g}{1+g} \left(\frac{1+g}{1+r} \right)^t \sum_{i=0}^s a_i \left(\frac{1+r}{1+r_f} \right)^{\min\{t,i\}} \quad (20)$$

Wagi a_i są tak dobrane, by wyeliminować wahania uznawane za charakterystyczne dla wahań koniunkturalnych, a więc takie, których cykl trwa krócej niż 20-36 kwartałów**. Zagadnienie doboru wag jest problemem z zakresu filtrowania. W omawianej pracy wybór wag odbywał się zgodnie ze zmodyfikowanym filtrem pasmowo-przepustowym (ang. *band-pass filter*)⁷.

W idealnym filtrze pasmowo-przepustowym trend ma postać dwustronnej średniej ruchomej:

$$\bar{C}_t^{id} = \sum_{i=-\infty}^{+\infty} \alpha_i \tilde{C}_{t-i} \quad (21)$$

gdzie wagi α_i dobierane są w ten sposób, aby z trendu wyeliminować wszystkie wahania o zadanych częstotliwościach. Narzędziem, które pozwala ocenić rozkład zmienności w danym szeregu według częstotliwości wahań, jest funkcja wzmocnienia (ang. *gain function*) $G(\omega)$, która pokazuje dla danej metody filtracji, jaka część fluktuacji o danej częstotliwości ω jest przez filtr przepuszczana. Dla idealnego filtra, którego celem jest eliminacja wpływu fluktuacji o częstotliwościach z określonego przedziału przyjmuje ona wartości 0 dla argumentów odpowiadającym wartościom z tego przedziału oraz 1 w pozostałych przypadkach.

W pracy do wyodrębnienia trendu zastosowano jednostronny przybliżony filtr pasmowo-przepustowy:

$$\bar{C}_t = \sum_{i=0}^m a_i \tilde{C}_{t-i} \quad (22)$$

* Wzór ten jest prawdziwy przy założeniu (7) oraz stałości stopy wolnej od ryzyka.

** Długość cyklu jest różna w różnych krajach. Generalnie w krajach rozwijających się cykle koniunkturalne są krótsze niż w krajach rozwiniętych.

⁷ Zob. M. Baxter, R. King, *Measuring Business Cycles: Approximate Band-Pass Filters for Economic Time Series*, NBER Working Paper, No. 5022, 1999; L. Christiano, T. Fitzgerald, *The Band-Pass Filter*, NBER Working Paper..., op. cit.; C. Schleicher, *Essays on Decomposition of Economic Variables*, 2003, www.iam.ubc.ca/theses/Schleicher/ChristophSchleicher_PhD_Thesis.pdf.

Przy czym wagi a_i dobierano teraz w taki sposób, by funkcja wzmocnienia filtra przybliżonego była jak najbardziej zbliżona do funkcji wzmocnienia filtra idealnego. Zagadnienie to sprowadza się więc do rozwiązania następującego problemu optymalizacyjnego:

$$\min_{a_i} \int_{-\pi}^{\pi} |G_a(\omega) - G_\alpha(\omega)|^2 f(\omega) d\omega \quad (23)$$

gdzie $f(\omega)$ jest funkcją gęstości spektralnej procesu poddawanego filtracji. Sposób rozwiązania powyższego problemu można znaleźć w pracach cytowanych na początku tej części⁸.

4. Wyniki

Koszty wahań badano dla kwartalnych danych polskich w okresie 3.1995-2.2010. Konsumpcja była reprezentowana przez dwie zmienne: wielkość spożycia ogółem oraz spożycie indywidualne. Dane dotyczące tych zmiennych pochodziły z rachunków narodowych publikowanych przez GUS. Kwartalne stopy wzrostu tych zmiennych równe były odpowiednio 0,97% oraz 1,05%.

W skład portfela aktywów bazowych wchodziły indeksy WIG20 oraz sWIG80, a także stopa wolna od ryzyka obliczana jako różnica pomiędzy średnim poziomem stopy procentowej WIBOR3m w danym kwartale a faktyczną stopą inflacji konsumenckiej w kwartale następnym. Średni poziom tak zdefiniowanej stopy wolnej od ryzyka wyniósł 1,31% na kwartał. Inflację konsumencką zastosowano również do obliczenia realnych stóp zwrotu z indeksów WIG20 oraz sWIG80.

Badając koszty wahań dla przełącznikowego modelu Markowa, wyróżniono dwa stany; fazę ożywienia, w której stopa wzrostu konsumpcji kształtowała się powyżej średniej oraz stagnacji, gdy stopa wzrostu konsumpcji była niższa niż średnia. Każda z faz musiała trwać co najmniej dwa okresy. Współczynniki opisujące trend dobrano tak, aby eliminowały wszystkie wahania, których okres jest krótszy niż 24 kwartały. Liczbę współczynników filtra ustalono na poziomie $m = 20$. Jako funkcję gęstości spektralnej analizowanego procesu przyjęto funkcję gęstości procesu AR(1) ze współczynnikiem autokorelacji równym 0,95.

⁸ M. Baxter, R. King, op. cit.; L. Christiano, T. Fitzgerald, op. cit.; C. Schleicher, op. cit.

Tabela 1

Koszty wszystkich wahań konsumpcji

Zmienna	r_f	g	Model	Koszt średni [%]	Koszt maks. [%]
Spożycie ogółem	1,31%	0,97%	IID	15,21 ($r = 1,36\%$)	146,07 ($r = 1,81\%$)
			MS	40,78 ($r = 1,45\%$)	182,91 ($r = 1,93\%$)
Spożycie indywidualne		1,05%	IID	22,27 ($r = 1,37\%$)	194,97 ($r = 1,81\%$)
			MS	56,51 ($r = 1,46\%$)	327,69 ($r = 2,16\%$)

Objaśnienia:

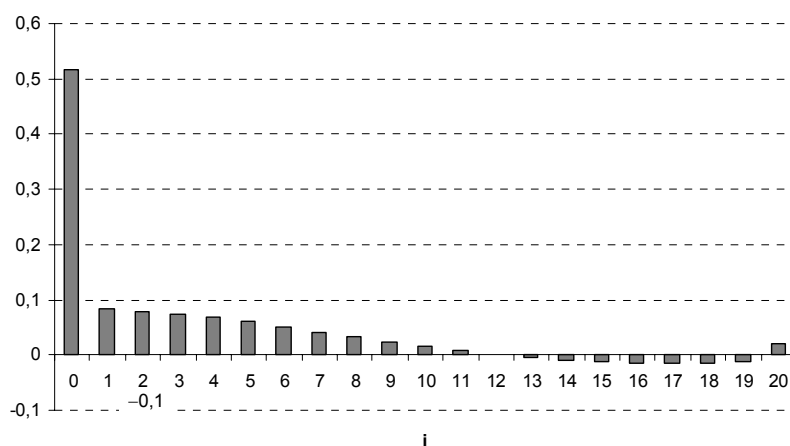
IID – model z niezależnymi wahaniami konsumpcji; MS – model przełącznikowy; r – stopa zwrotu z aktywa reprezentującego strumień aktualnej konsumpcji.

W pierwszej części badania wyznaczono koszty wszystkich wahań konsumpcji, korzystając ze wzorów (13) i (17) przy koszcie średnim oraz (15) i (18) dla kosztów maksymalnych. Uzyskane wyniki zestawiono w tabeli 1. Podano w niej także średnie stopy zwrotu z hipotetycznego aktywa, którego wypłaty replikują stopy wzrostu konsumpcji rzeczywistej.

Oszacowane średnie koszty wszystkich wahań konsumpcji zawierają się pomiędzy 15% a 60%. Oznacza to, że o taki procent powinna być zwiększona konsumpcja, aby użyteczność osiągnąca z takiego strumienia była równa użyteczności ze strumienia konsumpcji pozbawionej wszystkich fluktuacji. Wyniki te związane są z oszacowaniami stóp zwrotu z hipotetycznych aktywów reprezentujących rzeczywistą konsumpcję. Stopy te są z przedziału 1,36%-1,46%, a więc przewyższają stopę wolną od ryzyka o 0,05%-0,15%. Koszty maksymalne są zdecydowanie wyższe – zawierają się w przedziale 146%-328%, a odpowiadające im stopy r równe są 1,81%-2,16%. Maksymalna cena ryzyka konsumpcji $r - r_f$ waha się więc w granicach 0,50%-0,75% w skali kwartału. Wyniki są podobne dla obu zmiennych użytych do aproksymacji strumienia rzeczywistej konsumpcji. Ogólnie większe różnice występują pomiędzy modelem z niezależnymi stopami wzrostu konsumpcji (IID) a modelem przełącznikowym. W tym drugim przypadku koszty są mniej więcej dwukrotnie wyższe.

Wyniki te wskazują wyraźnie, że koszty wszystkich wahań konsumpcji są wysokie i jej wygładzenie wiązałoby się z dużymi korzyściami w kategoriach użyteczności ze strony gospodarstw domowych.

W drugiej części dokonano analizy kosztów wahań koniunkturalnych. W celu ich wyodrębnienia skorzystano z omówionego wcześniej filtra pasmowo-przepustowego. Na rysunku 1 zilustrowano wartości współczynników a_i ze wzoru (22) użyte do konstrukcji filtra.



Rys. 1. Oszacowane wartości współczynników a_i filtra pasmowo-przepustowego

W tabeli 2 zestawiono oszacowania kosztów wahań koniunkturalnych obliczone ze wzoru (20) oraz jego analogicznej wersji dla modelu przełącznikowego. Koszty średnie nie przekraczają wielkości 0,21%, natomiast koszt maksymalny jest nie większy niż 1,24% strumienia bieżącej konsumpcji. Wyniki te zdecydowanie różnią się od kosztów wszystkich wahań przedstawionych w poprzedniej tabeli. W porównaniu do nich są bardzo małe. Wynika stąd, że szczególnie kosztowne są wahania konsumpcji o niskich częstotliwościach, niższych niż wahania koniunkturalne. Dlatego uprawniona jest teza, że o ile wszelkie wahania konsumpcji mogą być rzeczywiście bardzo kosztowne dla gospodarstw domowych, o tyle dążenie do zmniejszania wahań koniunkturalnych nie przynosi istotnych korzyści. Należy mieć także na uwadze, że powyższe wyniki dotyczą kosztu krańcowego. Całkowite koszty wahań będą jeszcze niższe, nawet o połowę, w przypadku addytywnej funkcji oczekiwanej użyteczności.

Tabela 2

Koszty wahań koniunkturalnych konsumpcji

Zmienna	Model	Koszt średni [%]	Koszt maks. [%]
Spożycie ogółem	IID	0,08	0,73
	MS	0,20	0,91
Spożycie indywidualne	IID	0,08	0,74
	MS	0,21	1,24

Objaśnienie:

IID – model z niezależnymi wahaniami konsumpcji; MS – model przełącznikowy.

W pracy Alvareza i Jermanna średni koszt wszystkich wahań oscylował w granicach 20%-40%. Dość wysoki był koszt maksymalny, który przekraczał 200%, natomiast koszt wahań koniunkturalnych równy był około 0,1%, a maksymalny koszt nieznacznie przekraczał 0,5%. Oszacowana średnia cena ryzyka, mierzona jako kwartalna zyskowość papieru wartościowego uprawniającego do strumienia bieżącej konsumpcji, pomniejszona o stopę wolną od ryzyka, była równa około 0,1%. Maksymalna oszacowana wartość nie przekraczała 0,4%.

Wyniki uzyskane w pracy są więc dość zbliżone do oryginalnych rezultatów autorów. W ten sposób pokazano, że omawiana metoda daje podobne wyniki dla innych danych, nie pochodzących z rynku amerykańskiego. Tym bardziej, iż Alvarez i Jermann w swoich obliczeniach opierali się na danych rocznych, a w pracy wykorzystano dane kwartalne. Można więc uznać, że wyniki potwierdzają tezę, że prezentowana metoda nie jest specjalnie wrażliwa na sposób pomiaru kluczowych zmiennych.

Podsumowanie

W pracy przedstawiono nieparametryczną metodę oceny kosztów wahań konsumpcji zaproponowaną przez Alvareza i Jermanna. Wskazuje ona wyraźnie, iż mimo że wahania konsumpcji mogą być generalnie bardzo kosztowne dla gospodarstw domowych, to jednak zdecydowana większość tych kosztów związana jest z wahaniami długookresowymi o bardzo małej częstotliwości. Tym samym stosowanie kosztowych metod redukcji koniunkturalnych wahań w gospodarkach nie znajduje uzasadnienia.

Jednocześnie przeprowadzając analizę na danych polskich, pokazano, że uzyskane wyniki oraz wnioski nie różnią się istotnie od oryginalnych rezultatów autorów, dotyczących gospodarki amerykańskiej. W ten sposób potwierdzono tezę o względnej niewrażliwości metody na sposób pomiaru kluczowych zmiennych.

Omawiana metoda stanowi bardzo cenne narzędzie w dyskusji dotyczącej oceny korzyści płynących z działań na rzecz redukcji wahań cyklicznych w gospodarce. Jej główną zaletą jest nieparametryczność, to znaczy fakt, że nie opiera się ona na szczegółowych założeniach co do postaci funkcji użyteczności wykorzystywanej do oceny użyteczności płynącej ze strumienia konsumpcji.

Metoda jest również dobrym przykładem na możliwości wykorzystania informacji zawartych w cenach aktywów finansowych. Na podstawie obserwacji cen różnych klas aktywów pozwala na wnioskowanie o użyteczności strumienia konsumpcji wygładzonej w stosunku do konsumpcji rzeczywistej. Należy również zwrócić uwagę na fakt, że wprowadzanie nowych instrumentów finansowych będzie umożliwiało prowadzenie dokładniejszych badań w tym zakresie.

ASSET PRICES AND THE COST OF ECONOMIC FLUCTUATIONS IN POLAND

Summary

In the paper the method of measuring costs of business cycle fluctuations developed by Alvarez and Jermann (*Using Asset Prices to Measure the Cost of Business Cycles*, „Journal of Political Economy”, 2004, vol. 112, no. 6) is presented. The method is based mainly on asset prices data and doesn't need any assumptions regarding the utility function. Therefore it offers interesting alternative to parametric approaches. The results obtained for Polish data are consistent with Alvarez and Jermann's ones but reveal potential vulnerability to changes in the length of a measurement period, for example when switching from yearly to quarterly data.

Maria Balcerowicz-Szkutnik

PROBLEMY STARZENIA RYNKU PRACY – ANALIZA STATYSTYCZNO-DEMOGRAFICZNA DLA WYBRANYCH PAŃSTW UE

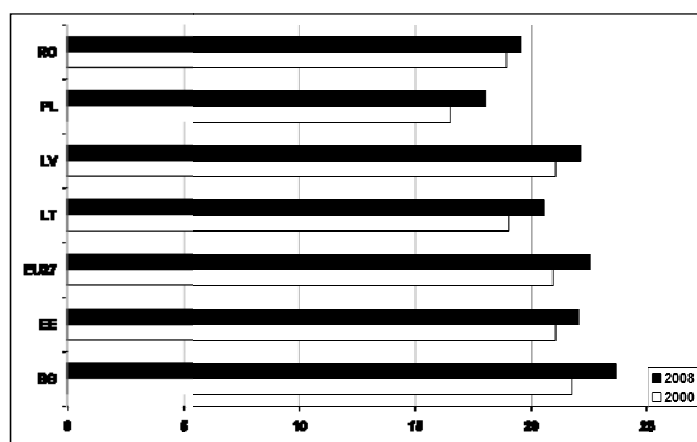
Wprowadzenie

Problemy zasobów pracy, siły roboczej i dostępności wykwalifikowanej kadry w jednostkach gospodarczych mają bezpośredni wpływ na funkcjonowanie tych jednostek w chwili obecnej i w przyszłości. Właściwa ocena aktualnego stanu rynku pracy oraz umiejętna i rzetelna ocena jego stanu przyszłego jest zatem nierozdzielnie związana z problematyką szeroko rozumianej strategii gospodarczej państw. Aktualny stan rynku pracy jest wypadkową wielu czynników, przy czym zdecydowanie najsilniejszy wpływ można przypisać zmianom demograficznym, a dokładniej zmiennej strukturze wieku ludności. W grupie osób w wieku produkcyjnym na szczególną uwagę zasługują osoby u schyłku kariery zawodowej, w wieku 50-64 lat, w literaturze przedmiotu nazywani pokoleniem 50+. Jest to specyficzna grupa pracowników, gdyż stanowią z reguły tzw. zaplecze zawodowe firmy, jako pracownicy o długoletnim doświadczeniu zawodowym, a z drugiej strony mogą stanowić swoiste obciążenie dla firm jako subpopulacja zagrożona największą absencją chorobową, niestabilnością zatrudnienia ze względu na możliwości korzystania ze świadczeń przedemerytalnych czy wreszcie najmniej skłonnością do rozszerzania swoich dotychczasowych kwalifikacji zawodowych i przez ten fakt nie zawsze pożądaną w przedsiębiorstwie. Jednocześnie jest to grupa pracowników w znacznym stopniu narażona na bezrobocie, przy czym w wielu przypadkach stan pozostawiania bez pracy jest niezależny od pracownika, lecz wymuszony zmienną strukturą gospodarczą. Problemy rynku pracy i dynamika jego parametrów dla pokolenia 50+ wydaje się zdecydowanie najciekawsza wśród wielu analiz związanych z aktywnością zawodową. Niniejsza tematyka będzie zatem przedmiotem prezentowanego artykułu. Ze względu na fakt, że na kształt rynku pracy ma obecnie znaczny wpływ wstąpienie Polski wraz z grupą innych państw do Unii Europejskiej, co umożliwiło swobodny przepływ siły roboczej wewnątrz UE,

analizy z zakresu zmienności rynku pracy będą obejmowały nie tylko Polskę, ale również inne państwa unijne. Ze względu na pewne tradycje i podobieństwa w funkcjonowaniu gospodarki wspólną analizą będą objęte następujące kraje: Estonia, Litwa, Łotwa, Bułgaria, Rumunia i Polska. Okres badań to lata 2000-2008, a bazę danych empirycznych zaczerpnięto z Eurostatu i urzędów statystycznych wymienionych państw.

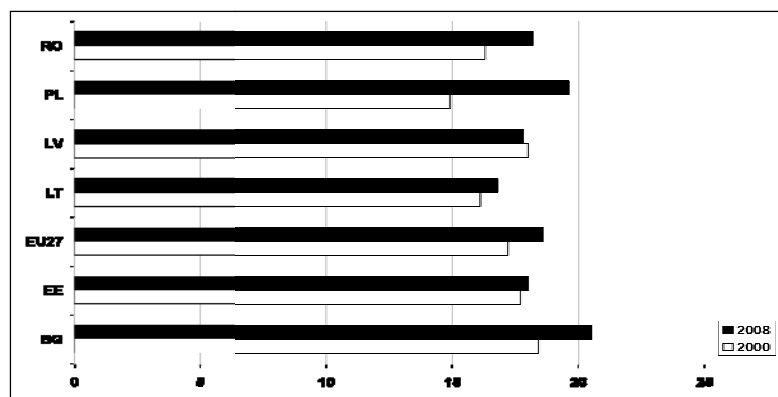
1. Zmienna dynamika struktury wieku pokolenia 50+

Jak wcześniej wspomniano, podstawowym czynnikiem kształtującym rynek pracy jest struktura wieku ludności. Jej zmienna dynamika wpływa w zdecydowany sposób na wartości parametrów szczegółowo opisujących rynek pracy, zatem dla wybranej subpopulacji osób przeprowadzimy wstępnie analizę jej zmian. Najwygodniej przedstawić dynamikę zmian na podstawie porównania stosownych piramid wieku ludności w dwóch wybranych latach. Analiza stosownych piramid wieku ludności zbudowanych dla lat 2000 i 2008 pozwala stwierdzić, że okres 2000-2008 charakteryzował się silnymi zmianami zarówno na szczycie piramidy, jak i u jej podstawy. Piramidy wieku dla 2008 r. w Polsce, Estonii oraz na Litwie i Łotwie mają formę wyraźne zwężoną u podstawy, co jest spowodowane drastycznym spadkiem liczby urodzeń w pierwszych latach XXI wieku. Najniższą dzietnością charakteryzowały się kobiety w Polsce i na Litwie (w 2008 r. 1,28 dziecka w przeliczeniu na 1 kobietę w wieku rozrodczym). Równocześnie obserwowany był intensywnie postępujący proces starzenia się społeczeństw. Społeczeństwo wkracza w fazę starości demograficznej, gdy odsetek ludzi w wieku 60 lat i więcej przekracza 12% ogółu społeczeństwa (zgodnie ze skalą starości zaproponowaną przez E. Rosseta). Dynamika zmian współczynnika starzenia przedstawiona na rysunku 1 pozwala na stwierdzenie, że społeczeństwo polskie oraz społeczeństwa rozważanych krajów UE wkroczyły w fazę starości demograficznej już znacznie wcześniej niż objęty analizą okres 2000-2008. Polska w porównaniu z pozostałymi państwami charakteryzuje się najniższą wartością miernika, czyli odsetek osób, które przekroczyły próg wieku emerytalnego w Polsce jest niższy niż w pozostałych państwach. W Bułgarii w 2008 r. już prawie co czwarta osoba była w wieku 60 lat i więcej, natomiast dwudziestoprocentowy próg wskaźnika został przekroczony w 2008 r. we wszystkich analizowanych krajach z wyjątkiem Polski i Rumunii.



Rys. 1. Udział osób w wieku 60 lat i więcej w wybranych krajach UE w latach 2000 i 2008

Jak wspomniano wcześniej, nazwą pokolenie 50+ określono grupę osób między 50 a 64 rokiem życia. Jak łatwo obliczyć, najstarsze osoby należące w 2008 r. do tej grupy urodziły się w 1944 r., czyli w czasie drugiej wojny światowej, najmłodsze natomiast – w 1958 r. W 2008 r. 18,6% (ponad 91 mln) ludności zamieszkującej 27 krajów Unii Europejskiej stanowiły osoby między 50 a 64 rokiem życia. W stosunku do 2000 r. udział tej grupy osób wzrósł o prawie 1,5%. Dynamikę zmian i geometryczny obraz tych współczynników przedstawia rysunek 2.



Rys. 2. Udział osób w wieku 50-64 lat w ogólnej liczbie ludności w wybranych krajach UE w latach 2000 i 2008

Źródło: Na podstawie danych Eurostatu.

Prognozy demograficzne głoszą, że w 2035 r. populacja osób w wieku 50-64 lat w UE wyniesie prawie 103 mln osób, co będzie stanowić około 20% ludności zamieszkującej obszar współczesnej UE. W tym samym czasie udział osób w wieku powyżej 65 lat sięgnie 25% ogólnej liczby ludności, zatem łącznie pokolenie osób powyżej 50. roku życia będzie liczyło w 2035 r. 235 mln osób, co będzie stanowić 45% wszystkich mieszkańców UE. Będzie to o prawie 10% więcej niż w 2008 r.

W perspektywie kilku najbliższych lat odsetek osób w wieku 50-64 lat będzie rósł we wszystkich analizowanych krajach. Tylko w Polsce w 2010 r. co piąta osoba będzie należała do populacji 50+. W Polsce natomiast spadek nastąpiłby po roku 2013. W tym okresie wiek 50-64 lat osiągną roczniki pochodzące z niżu demograficznego lat 60. XX w. Tylko w Estonii, na Litwie i Łotwie udział osób tej grupy wiekowej w ogólnej liczbie ludności nie przekroczy 19% w 2010 r. Jednakże po tym okresie prognozowany jest – w tych krajach – znaczny wzrost tego udziału, przy czym najwyższy na Litwie (2015 r. – 19,8%, 2020 r. – 21%).

2. Osoby w wieku 50+ na rynku pracy

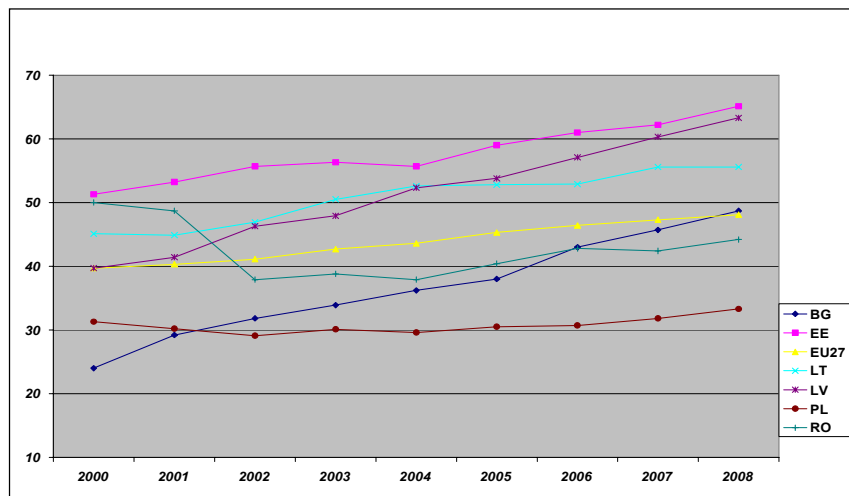
Ocenę stanu rynku pracy z uwzględnieniem poszczególnych grup wiekowych pracowników najlepiej przedstawić z wykorzystaniem podstawowych parametrów statystycznych, czyli wskaźników określających aktywność zawodową społeczeństwa i ewentualny poziom bezrobocia. Biorąc pod uwagę fakt, że Polska podobnie jak cała Europa, starzeje się i według długookresowych prognoz demograficznych w 2050 r. połowa ludności naszego kontynentu osiągnie wiek emerytalny, co spowoduje, że na unijnym rynku pracy zabraknie ponad 160 milionów pracowników, warto, jak wcześniej wspomniano, w prowadzonych analizach zwrócić szczególną uwagę na grupę wiekową pracowników narażonych na najsilniejszą dezaktywizację zawodową, czyli na osoby w wieku powyżej 50. roku życia.

Jak wykazano w pierwszej części opracowania, w najbliższych latach odsetek osób w wieku 50-64 lat będzie rósł i w związku z tym będą się nasilać problemy, takie jak: malejąca grupa osób aktywnych zawodowo, rosnące wydatki na świadczenia dla stanowiącej większość grupy osób biernych zawodowo czy pogłębiająca się społeczna marginalizacja osób starszych, co stanowi wielkie niebezpieczeństwo dla przyszłości systemów społeczno-ekonomicznych wielu krajów europejskich. Odrębnym problemem jest rosnąca stopa bezrobocia w tej grupie wiekowej pracowników, przy czym w wielu przypadkach jest to

bezrobocie strukturalne, czyli wymuszone niejako przez zmieniającą się gospodarkę i jej strukturę. Poniżej zostaną zaprezentowane wyniki szczegółowych analiz pozwalające na ocenę poziomu aktywności zawodowej i skali bezrobocia w grupie wybranych państw.

2.1. Aktywność zawodowa osób w wieku 50+

Dane bazy Eurostat pozwoliły na ocenę dynamiki zmian współczynników aktywności zawodowej za ostatnie lata. Interpretacja geometryczna empirycznych linii trendu współczynników dla ogółu pracowników przedstawiona na rysunku 3 pozwala na sformułowanie wielu ciekawych wniosków.



Rys. 3. Współczynniki aktywności zawodowej dla wybranych państw UE za lata 2000-2008

Poziom aktywności zawodowej różni się zdecydowanie w poszczególnych państwach. W państwach dawnego bloku radzieckiego współczynniki aktywności zawodowej oscylują w pobliżu poziomu ogólnounijnego, nieznacznie go przekraczając, natomiast w Polsce są zdecydowanie niższe niż w pozostałych krajach. Podobnie niższe wartości parametrów występują w przypadku Bułgarii i Rumunii. Niestety Polska jest jednym z państw UE o najniższym współczynniku aktywności zawodowej, zarówno ogółem, jak i w wyróżnionej grupie wiekowej 50+. Podobne badania prowadzone dla szerszej grupy państw wskazują na to, że niższy poziom współczynnika jest jedynie na Malcie.

Przyjmując założenie, że tempo wzrostu współczynnika aktywności zawodowej nie ulegnie zmianie, warto skonstruować prognozy jego wartości na najbliższe lata. Prognozy skonstruowano dwiema metodami – zarówno na podstawie funkcji trendu, jak i na podstawie metody wag harmoniczných, czyli metody adaptacyjnej, w której na wyniki prognozowania najsilniej wpływają nowo napływające informacje. Prognozy wyznaczone metodami adaptacyjnymi okazały się obarczone mniejszym błędem, zatem tylko takie zostaną zaprezentowane w artykule. Poniżej (tabela 1) przedstawiono wartości prognoz współczynników aktywności zawodowej w interesującej nas grupie wieku.

Tabela 1

Prognozy współczynników aktywności zawodowej w grupie wiekowej 50-64 lat ogółem dla wybranych państw UE na lata 2009-2011 (metoda wag harmoniczných)

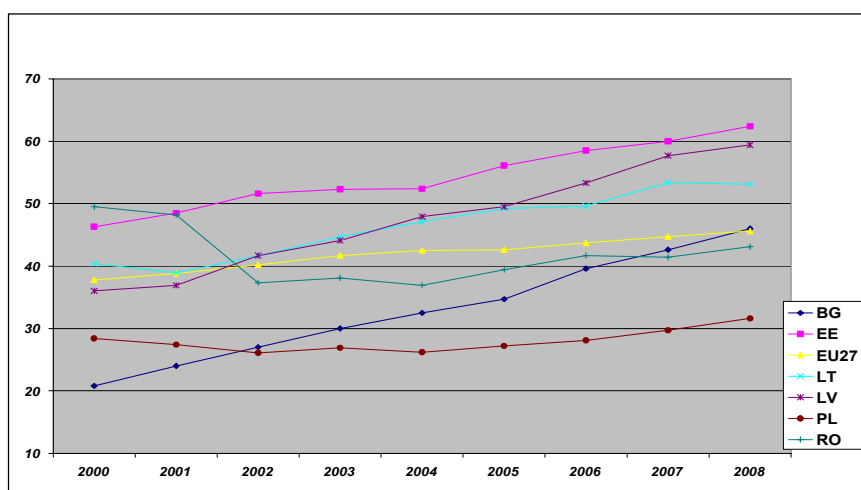
	BG	EE	EU27	LT	LV	PL	RO
2009 r.	51,64	66,70	49,14	57,31	66,35	34,05	44,42
2010 r.	54,63	68,58	50,16	58,57	69,37	34,86	45,00
2011 r.	57,61	70,46	51,18	59,83	72,38	35,67	45,58

Źródło: Obliczenia własne.

Wartości prognoz współczynnika potwierdzają wcześniej sformułowane spostrzeżenia, a mianowicie, że w krajach dawnych republik radzieckich, poziom aktywności zawodowej w grupie wiekowej 50-64 lat przekracza średni poziom unijny, natomiast w przypadku Polski aktywność zawodowa jest na poziomie niemal dwukrotnie niższym niż w Estonii, na Litwie czy Łotwie. Bułgaria i Rumunia to państwa niezbyt silnie odbiegające od poziomu UE pod względem aktywności zawodowej pokolenia 50+.

Obok poziomu aktywności zawodowej w analizie rynku pracy istotną rolę odgrywa miernik pozwalający na ocenę rzeczywistego natężenia siły roboczej, czyli wskaźnik zatrudnienia. Jest to miernik o bardzo silnej dynamice zmian i właśnie jego wartość pozwala na ocenę rzeczywistego poziomu aktywności zawodowej społeczeństwa, zwłaszcza w przypadku gdy ważna jest ocena aktywności zawodowej ludności w poszczególnych grupach wieku.

Na poniższym wykresie (rysunek 4) widoczna jest dynamika zmian wskaźników zatrudnienia w grupie wiekowej 50-64 lat w analizowanych państwach unijnych, a poniżej w tabeli 2 zamieszczono prognozy jego wartości na najbliższe lata wyznaczone metodą wag harmoniczných.



Rys. 4. Wskaźniki zatrudnienia dla wybranych państw UE za lata 2000-2008

Analizując empiryczne linie trendu wskaźnika zatrudnienia, zauważa się w zasadzie tę samą prawidłowość, która była widoczna w przypadku współczynników aktywności zawodowej, czyli zdecydowanie dominują pod względem wartości wskaźnika Litwa, Łotwa i Estonia. Bułgaria i Rumunia „doganiają” powoli przeciętny poziom w UE, natomiast w Polsce musi nastąpić jeszcze wiele zmian, najprawdopodobniej ustawowych, aby został osiągnięty przeciętny poziom Unii Europejskiej. W tabeli 2 zawarto prognozy wskaźnika aktywności zawodowej skonstruowane przy wcześniej przyjętych założeniach.

Tabela 2

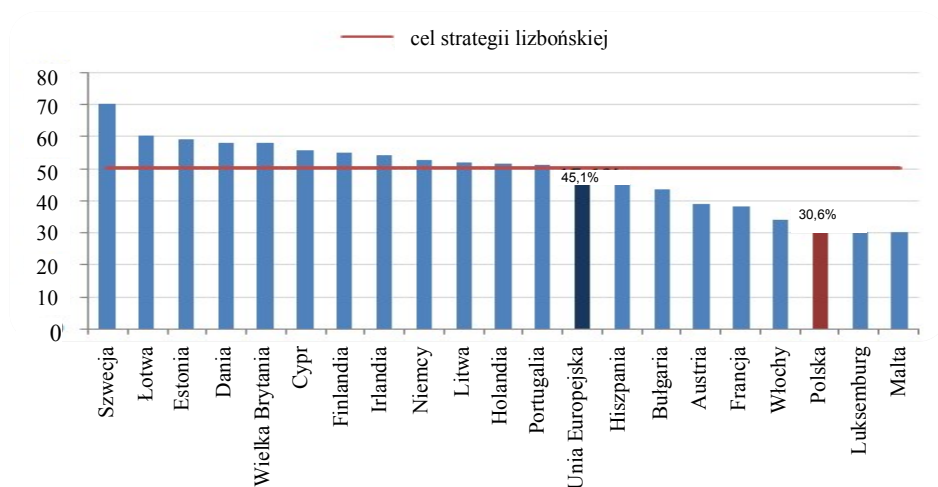
Prognozy wskaźnika zatrudnienia w grupie wiekowej 50-64 lat dla wybranych państw UE na lata 2009-2011 (metoda wag harmonicznych)

	BG	EE	EU27	LT	LV	PL	RO
2009 r.	49,19	64,27	46,53	55,54	62,94	32,70	43,31
2010 r.	52,46	66,30	47,44	57,30	66,03	33,85	43,85
2011 r.	55,72	68,32	48,36	59,06	69,12	35,00	44,38

Prognozy wskaźnika zatrudnienia potwierdzają wcześniejsze spostrzeżenia odnośnie do zmian współczynników opisujących poziom zatrudnienia na rynku pracy.

Problemy niskiej aktywności zawodowej i wręcz dezaktywizacji zawodowej pokolenia 50+ nękają nie tylko polski rynek pracy, lecz z takimi samymi problemami boryka się większość europejskich państw. Należało zatem podjąć próby znalezienia wspólnej formy ich rozwiązania. Taka próba została podjęta w marcu 2000 r. na szczycie Rady Europejskiej w Lizbonie, gdzie państwa członkowskie postawiły sobie za cel stworzenie z Unii Europejskiej najbardziej konkurencyjnej, opartej na wiedzy gospodarki na świecie. Cały proces, który miał prowadzić do wzrostu gospodarczego państw Unii, oraz co jest z tym związane, także do wzrostu zatrudnienia, został nazwany strategią lizbońską. Istotnym celem tej inicjatywy jest aktywizacja zawodowa osób z generacji 50+. Państwa Unii Europejskiej ustaliły, że do 2010 r. średni wiek przejścia na emeryturę wydłuży się z 58 do 63 lat, co jest przejawem próby aktywizacji zawodowej pokolenia 50+. Na szczycie w Lizbonie postanowiono również, że w tym samym czasie wskaźnik zatrudnienia dla osób powyżej 55. roku życia wzrośnie do 50%.

Realizacja zamierzeń strategii lizbońskiej spowodowała, że wskaźnik zatrudnienia w Unii Europejskiej w omawianej grupie wiekowej wynosił w IV kwartale 2007 r. (czyli po kilku latach od przyjęcia założeń lizbońskich) 45,1%. Zrealizować założenia z Lizbony udało się jedynie 12 z 27 państw członkowskich UE. Jak pokazuje rysunek 5, Polska pozostaje na jednym z ostatnich miejsc w tej klasyfikacji. Udział osób pracujących w ogólnej liczbie ludności w wieku 55-64 lat w naszym kraju, kształtował się na poziomie 30,6% w ostatnich miesiącach 2007 r., a w 2008 r. nieznacznie wzrósł.



Rys. 5. Wskaźnik zatrudnienia osób w wieku 55-64 lat w wybranych krajach UE w IV kwartale 2007 r. (dane za III kwartał 2007)

Źródło: Na podstawie danych Eurostatu.

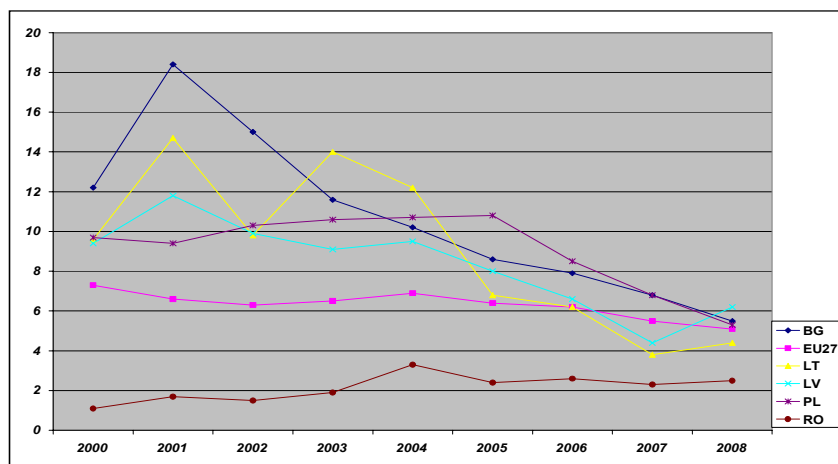
Niemal pewne jest, iż założenia z Lizbony z 2000 r., postulujące wzrost zatrudnienia osób w wieku 55-64 lat do poziomu 50%, pozostaną niezrealizowane w założonym czasie. Tylko nieliczne kraje mogą pochwalić się skuteczną polityką proaktywizacyjną i prozatrudnieniową skierowaną do najstarszych obywateli. Pozostali członkowie Wspólnoty, w tym Polska, nie zdążą z realizacją zaplanowanych celów do 2010 r. Według założeń Ministerstwa Pracy i Polityki Społecznej z lutego 2008 r. planuje się wzrost stopy osób pracujących w wieku 55-64 lat do poziomu 40% w 2010 r., a do poziomu 50% – dopiero w 2020 r.

Ogólnie można stwierdzić, że założenia strategii lizbońskiej są realizowane z mniejszym lub większym powodzeniem. Najwyższe wartości wskaźnika zatrudnienia przewiduje się w Szwecji, a najniższe niestety w Polsce. Bardzo dynamiczny wzrost w przypadku republik nadbałtyckich dawnego bloku radzieckiego wydaje się nieco zbyt optymistyczny, ale biorąc pod uwagę małą liczebność populacji tych państw, może być realny i możliwy do zrealizowania. W przypadku Polski, jak wcześniej wspomniano, nie należy się spodziewać gwałtownego wzrostu poziomu zatrudnienia w najbliższych latach. Decyzje dotyczące jego zmian i faktu aktywizacji zawodowej grupy wiekowej 50-64 lat mają charakter nie tylko gospodarczy, ale przede wszystkim polityczny i nie jest łatwo dokonać wyboru odpowiedniej strategii gospodarczej, tak aby były równocześnie uwzględnione racje ogólnospołeczne oraz polityczne.

2.2. Bezrobotni w wieku 50+ jako szczególna kategoria na rynku pracy

Przeciw wagą poziomu aktywności zawodowej jest poziom bezrobocia. Podobnie jak w przypadku poziomu aktywności zawodowej, ocena poziomu bezrobocia zostanie przeprowadzona na podstawie odpowiednich wskaźników statystycznych, czyli miernika stopy bezrobocia ogółem oraz z uwzględnieniem płci pracownika. Ze względu na fakt, że szczególnym problemem jest czas pozostawania pracownika bez pracy, odrębnymi badaniami będzie objęte tzw. bezrobocie długoterminowe.

Rysunek 6 przedstawia empiryczne linie trendu dla mierników stopy bezrobocia w latach 2000-2008 w rozważanych krajach UE.



Rys. 6. Dynamika zmian stopy bezrobocia ogółem za lata 2000-2008 w wybranych krajach UE

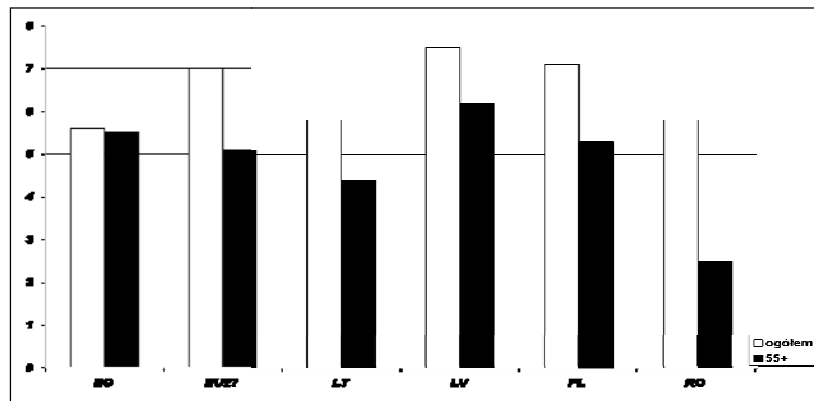
Źródło: Na podstawie danych Eurostatu.

Szczegółowa analiza powyższych linii trendu pozwala na następujące wnioski.

1. Dla ogółu państw UE stopa bezrobocia w ostatniej dekadzie uległa nieznacznemu spadkowi.
2. W latach 2000-2004 dynamika zmian współczynników była znacznie mniej intensywna niż w pozostałych latach analizowanego okresu i charakteryzowała się większą stabilnością, natomiast w latach 2005-2008 spadek stopy bezrobocia był zdecydowanie szybszy. Przyczyny tej zmiany można się dopatrywać w przypadku grupy wiekowej 50-64 lat przede wszystkim w zmianie statusu pracownika z bezrobotnego na przechodzącego w stan spoczynku zawodowego, czyli na wcześniejszą emeryturę. Liczba bezrobotnych nie obejmowała zatem osób, które według prawa nie były bezrobotne, czyli osób bez pracy, a pobierających zasiłek przedemerytalny, rentę lub wcześniejszą emeryturę.
3. Odrębne analizy prowadzone z uwzględnieniem płci pracownika potwierdzają wcześniej sformułowane wnioski, przy czym współczynniki określające stopę bezrobocia dla kobiet są niższe niż dla mężczyzn. Może to wynikać z faktu, że kobiety w większym stopniu korzystały ze świadczeń przedemerytalnych oferowanych przez zmieniony system przechodzenia na emeryturę. Warto przy tym zaznaczyć, że w latach 90. powszechny był pogląd, że dezaktywizacja zawodowa osób powyżej 50. roku życia jest zna-

komitym sposobem na zmniejszenie stopy bezrobocia. Powszechne było również przekonanie, że osoby te, przechodząc np. na wcześniejsze emerytury, zwalniają miejsca pracy dla bezrobotnych młodych ludzi. Co więcej, w Polsce taka właśnie polityka była realizowana, co zresztą jasno dowiodło jak bardzo błędny był to pogląd. Bezrobocie nie zmniejszyło się, bo zwalniane miejsca pracy nie były przeznaczone dla ludzi młodych i najczęściej były po prostu likwidowane. Zmniejszenie stopy bezrobocia w grupie wiekowej 50-64 lat było zatem pozorne i nie okazało się skutecznym lekarstwem na problemy rynku pracy.

O tym, że stopa bezrobocia w grupie wieku 50+ różni się od poziomu bezrobocia ogółem w całej zbiorowości pracowników, można się przekonać, porównując poziomy obydwu wskaźników w wybranym momencie czasu. Poniżej przedstawiono w formie graficznej obydwie współczynniki dla wybranych państw za 2008 r.



Rys. 7. Stopa bezrobocia w grupie wiekowej 50-64 lat w porównaniu z ogólnym poziomem bezrobocia dla państw nadbałtyckich w 2008 r. (ogółem)

Źródło: Na podstawie bazy Eurostatu.

Na rysunku widoczne jest, że poziom bezrobocia dla ogółu pracujących jest wyższy niż w grupie wiekowej 50-64 lat. Jest to w pewnym sensie zaskakujące, biorąc pod uwagę fakt, że pracodawcy w przypadku konieczności zwalniania pracownika kierują się kryterium dalszej jego przydatności, czyli spodziewanego dalszego czasu pracy. Jednak, jak wcześniej zasugerowano, jest to działanie pozorne, gdyż niekoniecznie na starych miejscach pracy pojawiają się nowi pracownicy, lecz są one po prostu likwidowane. Zatem po odejściu pracownika na wcześniejszą emeryturę stopa bezrobocia w grupie wiekowej

50+ maleje, a w przypadku ogółu pracowników pozostaje bez zmian lub wręcz rośnie. Dodatkowo na ogólny poziom bezrobocia ma wpływ również wysoka stopa bezrobocia wśród osób młodszych, czyli w grupie wiekowej 20-29 lat.

Biorąc pod uwagę silny związkowej stopy bezrobocia z rynkiem pracy, poziomem gospodarczym i społecznym oraz sprawnym funkcjonowaniem całego państwa, warto, podobnie jak w przypadku wcześniej rozpatrywanych mierników aktywności ekonomicznej ludności skonstruować prognozy stopy bezrobocia na najbliższe lata dla każdego z objętych analizą państw.

Tabela 3

Prognozy stopy bezrobocia w grupie wiekowej 50-64 lat dla wybranych państw UE na lata 2009-2011 (metoda wag harmoniczných)

	BG	EE	EU27	LT	LV	PL	RO
2009 r.	3,75	4,42	6,38	3,45	5,25	4,01	5,29
2010 r.	2,03	3,67	5,94	2,25	4,56	2,38	4,84
2011 r.	0,32	2,92	5,49	1,04	3,88	0,76	4,38

W wielu przypadkach prognozy uzyskane na 2011 r. wskazują na bardzo niski poziom stopy bezrobocia w grupie wiekowej 50-64 i są to prognozy raczej nierealne; należy się spodziewać stabilizacji jej poziomu lub niewielkiego spadku (o nieznacznej intensywności). Sytuacja obniżenia stopy bezrobocia do poziomu 0,67% (ogółem) w przypadku Polski jest absolutnie nierealna. Znajomość uwarunkowań gospodarczych i społecznych w kraju nie skłania do uznania tej prognozy za wiarygodną. Również w przypadku innych państw oszacowany poziom stopy bezrobocia nie może być uznany za wiarygodny, gdyż raczej nie są podejmowane zdecydowane kroki w zakresie polityki zatrudnienia mogące ograniczyć poziom bezrobocia. Ten spadek stopy bezrobocia, który obserwowano w ostatnich latach, był zapewne wynikiem pewnych nowych uwarunkowań prawnych w zakresie przyznawania większej liczby świadczeń przedemerytalnych, a nie (jak wspomniano wcześniej) gwałtownym wzrostem miejsc pracy, co teoretycznie powinno być przyczyną spadku stopy bezrobocia.

Podsumowanie

Zaprezentowane w artykule rozważania dotyczące wybranej grupy parametrów charakteryzujących rynek pracy miały na celu podkreślenie ważności problematyki niestabilności rynku pracy, nie tylko polskiego, ale i innych

państw unijnych. Oczywiście jest, że ze względu na rozmiary opracowania należało ograniczyć zakres analiz jedynie do kilku państw i wąskiej grupy parametrów charakteryzujących problem. Jak wykazano w powyższych analizach, polski rynek pracy w porównaniu z innymi państwami UE jest najbardziej obciążony problemami zarówno natury społecznej, jak i ekonomicznej.

Należy dodatkowo wspomnieć, że konsekwencje niestabilności zatrudnienia i wzrostu stopy bezrobocia pojawiają się i będą się pojawiać ze zwiększoną intensywnością w dziedzinach życia codziennego, w których dotychczas ich wpływ był nieznaczny i mało widoczny. Na przykład w sektorze bankowym, gdzie konsekwencje zmiennej dynamiki parametrów rynku pracy spowodują zanik płynności finansowej gospodarstw domowych, co spowoduje problemy ze ściągalsnością kredytów krótko- i długoterminowych, a zwłaszcza kredytów hipotecznych. Ponadto zupełnie naturalną konsekwencją niewypłacalności będzie zapewne zmiana zasad udzielania kredytów, co ograniczy ich dostępność dla przeciętnego pracownika. Okresowe zubożenie rodzin spowoduje, że zaniknie możliwość i chęć gromadzenia oszczędności w postaci lokat i innych instrumentów finansowych. Zatem problemy rynku pracy w przypadku ich nieefektywnego rozwiązania będą obejmować coraz szersze dziedziny gospodarki.

PROBLEMS OF AGING OF THE LABOR MARKET – STATISTICAL ANALYSIS-DEMOGRAPHIC FOR SELECTED EU COUNTRIES

Summary

This paper will present the results of detailed analysis of the basic parameters characterizing the dynamics of the labor market or the economic activity rate, employment rate, unemployment rate and long-term unemployment rate of workers aged 50-64 years for groups of countries newly joining the European Union. The analysis covers the next design trend function of these parameters is also an estimate of their value over the coming years designated several independent methods. A separate part of the development of an assessment of the impact parameters characterizing the labor market for basic economic and social processes.

Maciej Pichura

WPŁYW ASYMETRII ZMIENNOŚCI CEN NA EFEKTYWNOŚĆ INWESTYCJI NA RYNKU KAPITAŁOWYM

Wprowadzenie

Zmienność cen aktywów kapitałowych jest jednym z najistotniejszych czynników mających wpływ na decyzje podejmowane przez uczestników rynku. Ma ona znaczenie nie tylko w procesie inwestycyjnym, ale również w działaniach zabezpieczających inwestycje oraz podczas wyceny bardziej złożonych instrumentów rynkowych. Zmienność w tym obszarze jest bardzo często utożsamiana z ryzykiem, na które narażeni są uczestnicy rynku kapitałowego, co może prowadzić do stwierdzenia, że jest ona istotnym elementem większości działań na tym rynku. W tym artykule uwagę skoncentrowano jednak na aspekcie zależności efektów procesu inwestycyjnego w odniesieniu do zmienności.

Tradycyjne miary zmienności opisywane przez teorię finansów, takie jak wariancja, odchylenie standardowe czy też odchylenie przeciętne, charakteryzują się tym, że opisują rozproszenie w sposób symetryczny, jeśli chodzi o dodatnie i ujemne odchylenia od wartości przeciętnej. Zmienność wielu zjawisk w różnych obszarach badawczych mierzona w ten sposób jest zgodna z rzeczywistością jej interpretacją. W wypadku cen aktywów kapitałowych, gdy zmienność ma bardzo silne przełożenie na ryzyko, możliwe jest przyjęcie nieco odmiennych kryteriów jej określania. Aktywa kapitałowe są często wybierane jako potencjalny przedmiot inwestycji na podstawie analizy odchyleń ich stóp zwrotu od średniej w kierunku ujemnym, kierując mniej uwagi na odchylenia tego typu w kierunku dodatnim. Oczywiście istnieje jeszcze co najmniej kilka innych czynników, które mają wpływ na proces inwestycyjny, jednakże ich wpływ nie będzie tutaj szerzej omawiany. Koncepcja asymetrycznego charakteru zmienności stóp zwrotu istnieje w teorii finansów od ponad półwiecza¹ i wydaje się, że została uznana na wielu rynkach za faktycznie istniejącą. Szeroki zakres badań przeprowadzonych na różnych rynkach pozwala

¹ Zob. H. Markowitz, *Portfolio Selection. Efficient Diversification of Investment*, John Wiley & Sons, New York 1959; A.D. Roy, *Safety First and the Holding of Assets*, „Econometrica” 1952, Vol. 20, s. 431-449.

stwierdzić, że stopy zwrotu z przeważającej liczby aktywów² wykazują różnice w wartościach ze względu na wielkość odchyłeń dodatnich i ujemnych od średniej.

Jednym z możliwych zastosowań koncepcji asymetrii zmienności cen aktywów na rynkach kapitałowych może być analityczny pomiar zależności efektywności inwestycji ze względu na wrażliwość na zmienność rynkową wraz z jego aspektami prognostycznymi. Celem artykułu jest w związku z tym przedstawienie wybranych asymetrycznych miar zmienności, a także ukazanie możliwości ich zastosowania do analizy zależności efektywności inwestycji w odniesieniu do czynników wrażliwości wywodzących się od współczynnika beta zdefiniowanego w modelu CAPM. Analiza przedstawionych zagadnień zostanie przeprowadzona w odniesieniu do polskiego rynku akcji na podstawie rzeczywistych danych historycznych.

1. Zaproponowane miary zmienności

Podstawowe miary zmienności uwzględniające opisywany powyżej aspekt asymetrii to semiwariancja, semiodchylenie standardowe oraz semiodchylenie przeciętne. Są to miary, które w swoim założeniu jako „linię symetrii” przyjmują wartość średnią badanego zjawiska. Zatem, w zależności od intencji, wykazują one przeciętne odchylenia ujemne lub dodatnie wartości danego szeregu od jego średniej. Z drugiej strony można w łatwy sposób dostosowywać ich znaczenie informacyjne do preferencji poznawczych poprzez zastosowanie zamiast średniej jakiejś innej miary lub ustalonej stałej. W wypadku stóp zwrotu „linią symetrii” dla tych miar zmienności może być stopa zwrotu z rynku, stopa zwrotu wolna od ryzyka lub inna preferowana przez inwestora stopa zwrotu. Dla ustalenia uwagi nazwana ona zostanie wzorcową stopą zwrotu, a miary wyznaczane na jej podstawie będą miały przedrostek „skorygowany”.

Semiwariancję oraz semiodchylenie standardowe z uwzględnieniem stopy wzorcowej można ustalić, stosując następujące wzory³:

$$aSv^+ = \frac{1}{T} \sum_{i=1}^T \tau^+ (r_i - r_B)^2 \quad (1)$$

$$aSd^+ = \sqrt{aSv^+} \quad (2)$$

² Wyjątkiem może być rynek walutowy, jednakże niektóre badania [J. Wang, M. Yang, *Asymmetric Volatility in the Foreign Exchange Markets*, 2006] wykazują, że również on charakteryzuje się asymetryczną zmiennością, którą należy jednak analizować w nieco inny sposób.

³ Zob. H. Markowitz, op. cit.; V. Le Sourd, *Performance Measurement for Traditional Investment. Literature Survey*, EDHEC, Lille-Nice 2007.

$$aSv^- = \frac{1}{T} \sum_{i=1}^T \tau^-(r_i - r_B)^2 \quad (3)$$

$$aSd^- = \sqrt{aSv^-} \quad (4)$$

gdzie:

aSv^+ – skorygowana semiwariancja dodatnia stóp zwrotu z danego aktywu za okres T ,

aSd^+ – skorygowane semiodchylenie standardowe dodatnie stóp zwrotu z danego aktywu za okres T ,

aSv^- – skorygowana semiwariancja ujemna stóp zwrotu z danego aktywu za okres T ,

aSd^- – skorygowane semiodchylenie standardowe ujemne stóp zwrotu z danego aktywu za okres T ,

r_i – stopa zwrotu z aktywu w okresie i ,

$\tau^+ = 1$ dla $r_i > r_B$, $\tau^+ = 0$ dla $r_i \leq r_B$,

$\tau^- = 1$ dla $r_i \leq r_B$, $\tau^- = 0$ dla $r_i > r_B$,

r_B – wzorcowa stopa zwrotu (najczęściej średnia stopa zwrotu, ale może być to dowolna stopa określana według preferencji i oczekiwań inwestora).

Jeden z mierników współzależności statystycznej – kowariancja – jest konstruowany z wykorzystaniem miar rozproszenia opisanych powyżej. Podobnie jak semiwariancję czy semiodchylenie, można ją wyrazić, uwzględniając aspekt asymetrii zmienności. W takim wypadku, kierując się zastosowanymi już konwencjami nazewnictwa, skorygowana semikowariancja będzie miała następującą postać⁴:

$$aS\text{cov}^+(r_a, r_b) = \frac{1}{T} \sum_{i=1}^T \tau_a^+(r_{a,i} - r_{B,a}) \tau_b^+(r_{b,i} - r_{B,b}) \quad (5)$$

$$aS\text{cov}^-(r_a, r_b) = \frac{1}{T} \sum_{i=1}^T \tau_a^-(r_{a,i} - r_{B,a}) \tau_b^-(r_{b,i} - r_{B,b}) \quad (6)$$

gdzie:

$S\text{cov}^+(a, b)$ – skorygowana semikowariancja dodatnia stóp zwrotu z aktywów a oraz b za okres T ,

⁴ Zob. W. Cwynar, *Zmienność – dobra, czy zła? Analiza polskiego rynku kapitałowego*, „e-Finanse” 2010; Vol. 6, Nr 2; V. Le Sourd, op. cit.

$S\text{cov}^+(a,b)$ – skorygowana semikowariancja ujemna stóp zwrotu z aktywów a oraz b za okres T ,

$r_{a,i}$ – stopa zwrotu z aktywów a w okresie i ,

$r_{b,i}$ – stopa zwrotu z aktywów b w okresie i ,

$\tau_a^+ = 1$ dla $r_{a,i} > r_{B,a}$, $\tau_a^+ = 0$ dla $r_{a,i} \leq r_{B,a}$,

$\tau_a^- = 1$ dla $r_{a,i} \leq r_{B,a}$, $\tau_a^- = 0$ dla $r_{a,i} > r_{B,a}$,

$\tau_b^+ = 1$ dla $r_{b,i} > r_{B,b}$, $\tau_b^+ = 0$ dla $r_{b,i} \leq r_{B,b}$,

$\tau_b^- = 1$ dla $r_{b,i} \leq r_{B,b}$, $\tau_b^- = 0$ dla $r_{b,i} > r_{B,b}$,

$r_{B,a}$ – wzorcowa stopa zwrotu dla aktywów a (najczęściej jest to średnia stopa zwrotu z tego aktywów, ale można przyjąć inną stopę zwrotu),

$r_{B,b}$ – wzorcowa stopa zwrotu dla aktywów b (najczęściej jest to średnia stopa zwrotu z tego aktywów, ale można przyjąć inną stopę zwrotu).

Na podstawie przedstawionych powyżej mierników zmienności i współzależności możliwe jest wyznaczenie wzorów na współczynniki beta uwzględniające asymetrię zmienności stóp zwrotu. Beta to kluczowy parametr szeroko rozpowszechnionego modelu wyceny aktywów kapitałowych CAPM (ang. *Capital Asset Pricing Model*). Opisuje on relację ryzyka i zwrotu z danego aktywów (lub grupy aktywów, czyli portfela) w stosunku do rynku (najczęstszym przybliżeniem rynku jest główny indeks określonej giełdy). Współczynnik ten ukazuje wrażliwość zmian cen aktywów na ryzyko systematyczne (nie- dywersyfikowalne – wynikające z zachowań całego rynku). Jeśli beta jest równa zero – zmiany cen aktywów nie wykazują żadnego związku ze zmianami indeksu rynkowego. Interpretacja współczynnika beta ze względu na znak jest bardzo podobna do interpretacji współczynnika korelacji – dodatnie wartości pozwalają stwierdzić, że zmiany cen aktywów mają ten sam kierunek, co zmiany wartości indeksu rynkowego, a wartości ujemne – odwrotnie. Skala wartości tego współczynnika, która opisuje siłę zależności, nie ma natomiast ograniczeń. Beta indywidualnego aktywów lub portfela inwestycyjnego może być stosowana do oceny ich atrakcyjności inwestycyjnej. Pozwala on bowiem odnaleźć takie aktywów, których stopy zwrotu mają bardziej korzystne zmiany niż rynkowa stopa zwrotu. Współczynnik beta zdefiniowany w modelu CAPM ma następującą postać⁵:

$$\beta = \frac{\text{cov}(r, r_M)}{\sigma_M^2} \quad (7)$$

⁵ E.J. Elton et al., *Modern Portfolio Theory and Investment Analysis*, Wiley & Sons, New York 2002.

gdzie:

$\text{cov}(r, r_M) = \frac{1}{T} \sum_{i=1}^T (r_i - \bar{r})(r_{i,M} - \bar{r}_M)$ – kowariancja stóp zwrotu z aktywu ze wzorcową stopą zwrotu (z głównego indeksu danego rynku – w dalszych rozważaniach określona jako rynkowa stopa zwrotu) za okres T ,

σ_M^2 – wariancja rynkowej stopy zwrotu za okres T ,

$r_{i,M}$ – rynkowa stopa zwrotu w okresie i ,

r_i – stopa zwrotu z aktywu w okresie i ,

$\bar{r}_M$ – średnia rynkowa stopa zwrotu za okres T ,

$\bar{r}$ – średnia stopa zwrotu z aktywu za okres T .

Współczynniki beta uwzględniające asymetrię zmienności stóp zwrotu przyjmują następującą postać⁶:

$$a\beta^+ = \frac{aS \text{cov}^+(r, r_M)}{aSv^+(M)} \quad (8)$$

$$a\beta^- = \frac{aS \text{cov}^-(r, r_M)}{aSv^-(M)} \quad (9)$$

gdzie:

$a\beta^+$ – skorygowana „górna” beta aktywu,

$a\beta^-$ – skorygowana „dolna” beta aktywu,

pozostałe oznaczenia jak we wzorach (1)-(7).

Zaproponowany we wzorach (8) i (9) sposób pomiaru wrażliwości inwestycji zakłada, że wzorcowe stopy zwrotu zastosowane do obliczenia tych mierników to we wszystkich wypadkach średnia rynkowa stopa zwrotu z badanego okresu⁷.

Traktując opisane powyżej współczynniki beta jako miarę wrażliwości cen danego aktywu na zmienność rynkową, możliwe jest przeprowadzenie analizy efektywności inwestycji w zależności od wartości wspomnianych współczynników. W tym celu, obok zastosowania współczynników beta obliczanych za pomocą wzorów (7)-(9), jako kryteria analizy zostaną użyte

⁶ Za: A. Ang, J.S. Chen, Y. Xing, *Downside Risk*, „The Review of Financial Studies” 2006, Vol. 19, Iss. 4, s. 1191-1239.

⁷ Por. W. Cwynar, op. cit.

również następujące wyrażenia: $\beta - a\beta^-$, $\beta - a\beta^+$ oraz $a\beta^+ - a\beta^-$. Wyrażenia te oraz współczynniki beta zdefiniowane wcześniej zostaną w dalszej części artykułu określane bardziej ogólnie jako czynniki wrażliwości. Jednym z najbardziej interesujących zagadnień, które można przytoczyć w związku z opisywanymi czynnikami wrażliwości, jest kwestia znaczenia asymetrii zmienności oraz możliwości do uzyskania tzw. premii za ryzyko (czyli nadwyżki stopy zwrotu z inwestycji w aktyw ryzykowny nad stopę zwrotu z aktywa wolnego od ryzyka) w odniesieniu do rezultatów inwestycji kapitałowych. Zastosowanie zaproponowanych kryteriów analizy wydaje się zasadne z uwagi na fakt, że pozwoli ono na określenie ewentualnych dodatkowych walorów poznawczych w badaniu ryzykowności i efektywności inwestycji.

2. Metoda badawcza oraz wyniki przeprowadzonej analizy

Podstawowym celem przeprowadzonej analizy była próba określenia wpływu wrażliwości cen poszczególnych akcji na zmienność ogólną rynku na efektywnością inwestycji w te akcje ze szczególnym uwzględnieniem asymetrii zmienności. Badaniu został poddany polski rynek akcji w latach 2000-2010 z uwagi na fakt, że we wcześniejszym okresie istnienia GPW w Warszawie notowano na niej niedużą liczbę spółek oraz płynność obrotu była niezadowalająca. Badania prowadzono zarówno w pełnym okresie, jak i w dwóch podokresach – po arbitralnym wydzieleniu z głównego okresu części przypadającej na cykl wzrostowy (od 02/01/2003 do 30/06/2007) oraz spadkowy (01/07/2007 do 31/03/2009). Wszystkie dane zastosowanie w badaniu są rzeczywistymi notowaniami historycznymi pozyskanymi z serwisów www.parkiet.com oraz www.bossa.pl. W wypadku braku ciągłości danych zostały one uzupełnione na podstawie metody interpolacji liniowej. Jako wyznacznik zmian rynkowych został zastosowany główny indeks szerokiego rynku GPW – WIG.

Analiza zależności efektywności inwestycyjnej spółek giełdowych od opisanych czynników wrażliwości została podzielona na kilka etapów. W pierwszej kolejności akcje zostały uszeregowane rosnąco ze względu na jeden z czynników wrażliwości. Na podstawie tego uszeregowania utworzono siedem równolicznych grup akcji. Dla każdego z trzech badanych okresów przeprowadzono zatem grupowanie spółek na podstawie następujących współczynników: β , $a\beta^+$, $a\beta^-$ oraz wyrażen: $\beta - a\beta^-$, $\beta - a\beta^+$, $a\beta^+ - a\beta^-$. Następnie dla każdego przedziału obliczona zostaje średnia roczna stopa zwrotu z inwestycji we wszystkie spółki do niego należące. Jak można zauważyć, jako miara efektywności inwestycji została wybrana stopa zwrotu. Jest to najprostsza miara oceniająca rezultaty inwestycji i nie do końca można ją utożsamiać stricte z efektywnością, której miara powinna uwzględniać jeszcze inne, oprócz

zwrotu, aspekty inwestycji. W niniejszych rozważaniach analizie jest poddawana zależność zyskowności prostej inwestycji typu „kup i trzymaj” od wrażliwości na zmienność rynkową wyrażaną na kilka sposobów, na podstawie współczynnika beta. Z uwagi na fakt, że współczynnik beta niesie ze sobą informacje, które wpływają na określenie efektywności inwestycji, nie wydaje się konieczne zastosowanie bardziej złożonych od stopy zwrotu miar oceniających tę efektywność. Dla każdego z wyznaczonych przedziałów obliczane są również średnie wartości współczynników β , $a\beta^+$ oraz $a\beta^-$, co pozwala na uzyskanie dodatkowych informacji możliwych do wykorzystania w celach porównawczych.

Wyniki przeprowadzonych analiz zostały przedstawione w tabelach 1-3, które podzielono ze względu na badany okres. W tabeli 1 zaprezentowano rezultaty badania pełnego okresu (lata 2000-2010). Analizując wyniki badania za pełny okres przyjęty we wcześniejszych założeniach, należy zwrócić uwagę na kilka wniosków. Pierwszy z nich dotyczy faktu, iż grupowanie aktywów z zastosowaniem tradycyjnego współczynnika β daje w efekcie w zasadzie tylko jedną grupę spółek odznaczających się bardzo wysoką średnią roczną stopą zwrotu. Jest to jedyny przypadek, gdy dodatnie zmiany cen akcji są silniejsze niż rynkowe. Metody grupowania z zastosowaniem $a\beta^+$ nie pozwala na jednoznaczne wskazanie grupy o najwyższej średniej rocznej stopie zwrotu. Natomiast w wypadku zastosowania grupowania według wartości wyrażenia $\beta - a\beta^+$, można wyróżnić grupę z bardzo wysoką średnią stopą zwrotu, jednakże cztery z pozostałych sześciu grup spółek charakteryzują się stopą zwrotu na porównywalnym poziomie, co daje małe możliwości różnicowania. Stosując metody grupowania opierające się o współczynniki $a\beta^-$, można wskazać w sposób najbardziej jednoznaczny grupy spółek charakteryzujące się bardzo wysokimi średnimi rocznymi stopami zwrotu. Co więcej, łatwiej jest również w tym wypadku wyróżnić takie spółki, które wykazują średnie roczne stopy zwrotu zdecydowanie niższe niż pozostałe. Można także zauważyć, że średnie wartości skorygowanej „dolnej” bety w sytuacjach, gdy zostały wyodrębnione grupy spółek o ponadprzeciętnie wysokich stopach zwrotu, są bardzo wysokie. Istnienie takiej zależności może sugerować, że osiąganie wysokich stóp zwrotu w pewien sposób wiąże się z narażeniem kapitału na zdecydowanie większe ryzyko. Wysoka wartość współczynnika $a\beta^-$ jest bowiem skutkiem znacznie większej zmienności cen spółek w niekorzystnym kierunku w porównaniu do zmian cen ogólnorynkowych w tym samym okresie.

Tabela 1

Wyniki analizy zależności efektów inwestycji od czynników wrażliwości
w latach 2000-2010

Grupowanie według β					Grupowanie według $\beta - a\beta^+$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	35,47%	-2,65	4,88	7,33	1	16,48%	-0,66	7,13	7,80
2	3,76%	0,15	1,02	1,82	2	242,32%	0,44	2,12	3,81
3	-4,62%	0,26	1,08	1,94	3	-6,51%	0,42	1,47	2,22
4	2,35%	0,40	1,23	2,26	4	18,43%	0,56	1,34	2,02
5	7,75%	0,54	1,30	1,67	5	15,88%	0,60	1,19	1,71
6	12,45%	0,70	1,85	2,21	6	24,54%	0,61	0,99	1,50
7	236,22%	1,55	3,50	4,05	7	-20,12%	-1,05	0,35	2,05
Grupowanie według $a\beta^+$					Grupowanie według $\beta - a\beta^-$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	1,80%	0,27	1,57	0,16	1	337,99%	-0,45	3,66	10,41
2	103,51%	0,89	1,48	0,34	2	-4,38%	0,25	3,43	3,07
3	13,23%	1,12	2,23	0,43	3	6,01%	0,39	2,09	2,05
4	9,39%	1,32	1,73	0,55	4	-29,14%	0,38	1,09	1,60
5	8,60%	1,62	2,46	0,57	5	9,75%	0,44	1,10	1,38
6	20,94%	2,15	4,41	0,50	6	-4,44%	0,59	1,29	1,29
7	135,03%	7,74	7,28	-1,63	7	-25,09%	-0,68	2,04	1,08
Grupowanie według $a\beta^-$					Grupowanie według $a\beta^+ - a\beta^-$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	-111,80%	0,36	2,78	0,89	1	393,64%	-0,23	2,67	10,13
2	8,45%	0,42	0,99	1,21	2	4,42%	0,27	0,99	2,44
3	1,78%	0,45	1,63	1,44	3	-3,18%	0,36	1,30	2,15
4	-24,52%	0,49	1,72	1,68	4	16,57%	0,49	1,25	1,64
5	0,84%	0,44	1,57	2,10	5	13,62%	0,51	1,18	1,36
6	11,81%	0,36	2,57	3,22	6	-21,75%	0,61	1,44	1,40
7	408,49%	-1,61	3,54	11,04	7	113,62%	-1,09	6,16	1,80

Objaśnienia:

$\bar{r}$ – średnia roczna stopa zwrotu ze wszystkich spółek w danej klasie,

$\bar{\beta}$ – średnia wartość współczynnika β w danej klasie,

$\overline{a\beta^+}$ – średnia wartość współczynnika $a\beta^+$ w danej klasie,

$\overline{a\beta^-}$ – średnia wartość współczynnika $a\beta^-$ w danej klasie.

Średnia roczna rynkowa stopa zwrotu dla tego okresu wyniosła 13,23%.

Konkludując, w pełnym badanym okresie można zauważyć, że najbardziej efektywne pod względem średniej rocznej stopy zwrotu były akcje, których ceny wykazywały silną wrażliwość na zmienność rynkową. Dodatkowo można stwierdzić, że największą wartość wyjaśniającą w tej analizie wykazują czynniki wrażliwości uzależnione od skorygowanej „dolnej” bety (w niektórych przypadkach było możliwe wyodrębnienie spółek o stopach zwrotu przewyższających średnią rynkową stopę zwrotu z tego okresu ponad 25-krotnie).

W tabeli 2 zostały zaprezentowane wyniki badań zależności stóp zwrotu od czynników wrażliwości dla okresu, który został określony we wcześniejszych założeniach jako hossa. W tym wypadku zależność pomiędzy wrażliwością cen akcji na zmienność rynkową a wynikami inwestycji jest bardzo wyraźna. W większości przypadków im silniejsza ta wrażliwość, tym wyższe są wykazywane przez akcje stopy zwrotu. Grupowanie według współczynnika β pozwala w pewnym stopniu odnaleźć spółki, które wykazują niższe stopy zwrotu niż pozostałe, jednak spośród pozostałych pięciu grup spółek, które charakteryzują się bardzo dobrymi rezultatami inwestycyjnymi, trudno jest wyróżnić takie, których średnia roczna stopa zwrotu jest znacząco wyższa od pozostałych. W wypadku zastosowania w grupowaniu czynników $a\beta^+$ oraz $\beta - a\beta^+$ jeszcze trudniejsze jest zróżnicowanie spółek ze względu na efektywność inwestycyjną. Z drugiej strony pozwala to na uzyskanie kilku różnych grup spółek wykazujących wysokie średnie roczne stopy zwrotu i skupić dalsze analizy na innych aspektach oczekiwanych przez inwestora.

Tabela 2

Wyniki analizy zależności efektów inwestycji od czynników wrażliwości
w okresie określonym jako cykl wzrostowy

Grupowanie według β					Grupowanie według $\beta - a\beta^+$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	112,18%	0,13	1,73	2,38	1	305,78%	0,54	3,62	4,58
2	72,66%	0,43	1,87	1,71	2	291,25%	0,68	1,80	2,70
3	270,63%	0,52	1,20	1,74	3	188,33%	0,67	1,55	2,07
4	232,97%	0,61	1,40	2,01	4	178,06%	0,59	1,30	1,82
5	333,17%	0,73	1,70	3,06	5	364,56%	0,63	1,22	1,79
6	256,79%	0,86	1,69	2,35	6	146,42%	0,59	1,05	2,22
7	316,10%	1,25	2,07	3,36	7	115,70%	0,81	1,08	1,50

cd. tabeli 2

Grupowanie według $a\beta^+$					Grupowanie według $\beta - a\beta^-$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	101,66%	0,39	0,82	1,27	1	482,00%	0,75	2,89	5,66
2	169,18%	0,51	1,05	1,40	2	279,77%	0,68	1,74	2,80
3	181,78%	0,62	1,20	2,42	3	306,18%	0,63	1,50	2,11
4	282,14%	0,74	1,36	2,01	4	273,66%	0,59	1,33	1,81
5	215,68%	0,63	1,58	2,25	5	166,24%	0,60	1,19	1,58
6	353,09%	0,84	1,89	2,61	6	86,59%	0,59	1,12	1,35
7	291,21%	0,80	3,77	4,78	7	-8,34%	0,68	1,87	1,16
Grupowanie według $a\beta^-$					Grupowanie według $a\beta^+ - a\beta^-$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	47,08%	0,45	1,01	1,04	1	569,84%	0,84	2,26	5,38
2	97,87%	0,58	1,06	1,36	2	259,87%	0,62	1,80	2,81
3	138,48%	0,61	1,20	1,58	3	406,98%	0,62	1,29	1,99
4	264,29%	0,68	1,37	1,87	4	142,27%	0,62	1,29	1,79
5	328,50%	0,73	1,61	2,23	5	163,59%	0,62	1,28	1,61
6	307,42%	0,70	1,83	2,88	6	66,83%	0,64	1,17	1,34
7	412,35%	0,76	3,60	5,94	7	-25,28%	0,55	2,56	1,57

* Oznaczenia jak w tabeli 1. Średnia roczna rynkowa stopa zwrotu dla tego okresu wyniosła 80,04%.

Zastosowanie metod grupowania zależnych od współczynnika $a\beta^-$ ponownie daje najlepszy efekt pod względem różnicowania spółek, jeśli chodzi o średnią roczną stopę zwrotu. Szczególnie widoczne jest to w przypadku grupowania przy użyciu wyrażeń $\beta - a\beta^-$ i $a\beta^+ - a\beta^-$, gdyż jego zastosowanie jako jedyne pozwala na wyodrębnienie grup spółek wykazujących ujemne średnie roczne stopy zwrotu. Taki wynik daje dużą wartość poznawczą zważywszy na fakt, że badany okres to hossa, a pozostałe metody grupowania pozwoliły na wyselekcjonowanie tylko trzech grup spółek, których średnia roczna stopa zwrotu była niższa niż 100% (47% i 98% dla $a\beta^-$ oraz niecałe 73% w wypadku zastosowania tradycyjnej β jako kryterium grupującego). Dla porównania, średnia roczna stopa zwrotu z indeksu WIG wyniosła w tym okresie nieco ponad 80%. Podobnie jak w przypadku analizy pełnego okresu badawczego, wartości współczynnika $a\beta^-$ dla grup spółek

o zdecydowanie ponadprzeciętnych średnich stopach zwrotu są wysokie. Wydaje się to potwierdzać wcześniejsze przypuszczenie, że osiąganie wysokich stóp zwrotu jest powiązane ze zdecydowanie większym zagrożeniem zmianą cen w negatywnym kierunku.

Wyniki analizy badanej zależności dla okresu, który został określony w założeniach początkowych jako *bessa*, zostały przedstawione w tabeli 3. Prawie w każdym wariancie grupowania spółek można wyodrębnić jedną grupę, która w zdecydowany sposób wykazuje niskie średnie roczne stopy zwrotu. Jedynie zastosowanie do klasyfikacji wyrażenia $\beta - a\beta^+$ skutkuje wydzieleniem dwóch skrajnych grup o najniższych stopach zwrotu. Z drugiej strony, prawie wszystkie kryteria grupowania, oprócz wyróżnienia spółek o bardzo niskich stopach zwrotu, nie pozwalają na większe zróżnicowanie ze względu na rezultaty inwestycyjne pozostałych grup. W przypadku zastosowania współczynnika $a\beta^+$ możliwe było wydzielenie grupy spółek o względnie wysokich ($-8,52\%$), ale mimo wszystko ujemnych średnich stopach zwrotu. Jedynie grupowanie ze względu na $a\beta^+ - a\beta^-$ dało możliwość wyodrębnienia grupy spółek, których średnia roczna stopa zwrotu była dodatnia (co więcej, wynosiła ona prawie $17,5\%$ przy średniej rynkowej stopie zwrotu na poziomie $-36,2\%$).

Tabela 3

Wyniki analizy zależności efektów inwestycji od czynników wrażliwości
w okresie określonym jako cykl spadkowy

Grupowanie według β					Grupowanie według $\beta - a\beta^+$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	-166,36%	-1,50	13,23	1,72	1	-112,73%	0,13	15,14	2,70
2	-44,73%	0,33	1,70	1,43	2	-49,75%	0,67	2,07	2,10
3	-48,84%	0,50	1,36	1,23	3	-44,46%	0,62	1,49	1,31
4	-39,59%	0,62	1,56	1,27	4	-22,84%	0,66	1,34	1,25
5	-45,15%	0,77	1,74	1,50	5	-39,88%	0,76	1,31	1,38
6	-40,81%	0,97	1,61	1,55	6	-35,35%	0,74	1,17	1,21
7	-41,76%	1,58	3,16	2,52	7	-123,24%	-0,35	1,17	1,24

cd. tabeli 3

Grupowanie według $a\beta^+$					Grupowanie według $\beta - a\beta^-$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	-8,52%	0,30	0,79	1,10	1	-34,94%	0,39	14,11	3,82
2	-43,26%	0,56	1,11	1,11	2	-86,95%	0,59	2,20	1,66
3	-44,26%	0,63	1,28	1,22	3	-55,01%	0,55	1,47	1,26
4	-42,80%	0,73	1,42	1,29	4	-40,74%	0,57	1,30	1,13
5	-40,76%	0,92	1,65	1,61	5	-38,90%	0,71	1,32	1,15
6	-27,27%	0,76	2,10	2,00	6	-39,46%	0,70	1,25	1,04
7	-222,57%	-0,66	16,35	2,91	7	-132,42%	-0,28	2,09	1,12
Grupowanie według $a\beta^-$					Grupowanie według $a\beta^+ - a\beta^-$				
Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$	Klasa	$\bar{r}$	$\bar{\beta}$	$\overline{a\beta^+}$	$\overline{a\beta^-}$
1	-125,31%	0,14	11,35	0,61	1	17,44%	0,61	1,93	3,32
2	-29,51%	0,49	1,31	0,97	2	-43,11%	0,74	1,31	1,43
3	-45,90%	0,57	1,39	1,12	3	-39,33%	0,73	1,27	1,25
4	-44,46%	0,72	1,47	1,27	4	-49,67%	0,71	1,32	1,21
5	-59,92%	0,82	2,44	1,46	5	-41,64%	0,68	1,43	1,18
6	-81,79%	0,94	2,44	1,77	6	-36,14%	0,65	1,66	1,19
7	-40,87%	-0,44	3,78	4,07	7	-237,30%	-0,89	15,76	1,62

* Oznaczenia jak w tabeli 1. Średnia roczna rynkowa stopa zwrotu dla tego okresu wyniosła -36,2%.

Badając okres bessy, nie można już wyraźnie zauważyć zależności pomiędzy ponadprzeciętnymi wartościami stóp zwrotu a dużą wrażliwością spółek na negatywne zmiany rynkowe (o czym świadczyła poprzednio wysoka wartość współczynnika $a\beta^-$). Mimo to dla jedynej grupy spółek wykazującej dodatnie stopy zwrotu zależność ta wydaje się być silna. Interesujący natomiast jest fakt, iż we wszystkich przypadkach grupowania, spółki mające najniższe średnie roczne stopy zwrotu mają bardzo wysoką wartość $\overline{a\beta^+}$. Jest to trudne do wytłumaczenia tym bardziej, że we wszystkich tych grupach wartość średnia współczynnika beta jest ujemna. Wskazuje to na istnienie bardzo silnej wrażliwości tych spółek na pozytywną zmienność rynkową. Bardziej wnikliwa analiza tych przypadków pozwoliła natomiast na wytłumaczenie tego faktu. Wysoka wartość $\overline{a\beta^+}$ oraz ujemna wartość $\bar{\beta}$ wynikają z kilku bardzo skrajnych dodatnio obserwacji, które zaburzyły te wartości.

Podsumowanie

W wielu opracowaniach dotyczących zmienności cen na rynkach kapitałowych⁸ można znaleźć potwierdzenie przekonania, że ma ona charakter asymetryczny. Ponadto, niektóre z nich pozwalają na wysunięcie wniosków, że ta asymetria ma duży wpływ na możliwe do uzyskania rezultaty inwestycji. Istnieją również stwierdzenia, że aktywa, które wykazują większą wrażliwość na niekorzystne zmiany rynkowe, pozwalają na uzyskanie wyższych stóp zwrotu. Analiza przedstawiona w tym artykule miała na celu próbę określenia, czy i w jakim stopniu zależności o podobnym charakterze istniały na polskim rynku akcji w okresie od początku 2000 r. do końca trzeciego kwartału 2010 r. Dodatkowo, zostały przedstawione pokrótce mierniki (nazwane czynnikami wrażliwości), które zastosowano do przeprowadzenia badania.

Wyniki przeprowadzonej analizy dają podstawy do potwierdzenia zgodności rezultatów badań na rynku polskim w odniesieniu do wcześniejszych badań przeprowadzonych na międzynarodowych rynkach. Można bowiem dość wyraźnie zaobserwować dwa ciekawe aspekty zastosowania skorygowanych współczynników beta jako czynników tłumaczących wysokość średnich stóp zwrotu z polskich akcji. Pierwszy z nich wskazuje na co najmniej dobre właściwości grupujące skorygowanego „dolnego” współczynnika beta. Zatem możliwe jest wysunięcie przypuszczenia, że stopy zwrotu są w pewien sposób uzależnione od wrażliwości zmian cen na zmienność rynkową z naciskiem na uwzględnienie zmian w kierunku negatywnym. Drugi aspekt w pewnym sensie wynika z pierwszego. Mianowicie, jeśli założyć istnienie powyżej opisanej zależności, należy również stwierdzić, że wysoka wrażliwość zmian cen na zmiany rynkowe ma przełożenie na możliwość uzyskania ponadprzeciętnej stopy zwrotu, co można w pewien sposób uznać za premię za ryzyko.

Zaprezentowana metoda badawcza pozwala na szczegółowe wykonanie analizy zależności istniejących między efektywnością inwestycji a jej wrażliwością na zmienność rynkową. Dodatkowo można przypuszczać, iż znajdzie ona również aplikację w teorii konstrukcji portfeli inwestycyjnych. Innym zastosowaniem, które nie zostało tutaj opisane, a może zostać wykorzystane, jest aspekt prognostyczny przedstawionych współczynników, który wymaga jednak szerszego zbadania. Możliwe jest wreszcie zastosowanie zaproponowanych miar wrażliwości w konstrukcji analitycznych mierników i metod szeroko pojętej oceny efektywności inwestycyjnej, gdzie własności rynku związane z asymetrią są w ostatnim czasie coraz częściej brane pod uwagę.

⁸ Między innymi: A. Ang, J.S. Chen, Y. Xing, op. cit.; G. Bekaert, G. Wu, *Symmetric vs. Downside Risk: Does It Matter Portfolio Choice?*, „The Review of Financial Studies” 2000, Vol. 13, No. 1; J. Wang, M. Yang, op. cit.

ASYMMETRIC VOLATILITY MEASURES AND INVESTMENT EFFICIENCY – CHOSEN APPLICATIONS IN CAPITAL MARKETS

Summary

This article shows chosen issues concerning volatility asymmetry on capital markets. In detail, it turns attention on Polish stock exchange in period between 2000 and 2010, which is additionally divided into rising and descending market periods. Asymmetric volatility measures have been used to describe adjusted beta coefficients. These coefficients have been named sensitivity factors and they have been used to classify stocks regarding yearly average rate of return.

Analysis of stock return dependency on their prices sensitivity to market volatility may result in assuming that presented coefficients give chance, in most cases, to find groups of stocks which prove to have extraordinary rates of return. Beta coefficient based on downside volatility measures seems to have biggest cognitive value. It is boldly visible especially when analyzing down-trending market. The results that have been obtained in this research may also lead to assumption that stocks being very sensitive to negative market volatility tend to give so called risk premium.

Wanda Ronka-Chmielowiec*

UBEZPIECZENIA KOMUNIKACYJNE JAKO PODSTAWOWE PRODUKTY UBEZPIECZENIOWE SPRZEDAWANE NA POLSKIM RYNKU UBEZPIECZENIOWYM

Wprowadzenie

Tematyka przedstawionego artykułu jest związana z funkcjonowaniem ubezpieczeń komunikacyjnych w Polsce. Rynek ubezpieczeniowy po transformacjach systemu gospodarczego działa już od 20 lat, jednak ciągle jest w fazie rozwoju, pomimo że przeszedł już wiele zmian. Ubezpieczenia komunikacyjne ciągle mają największy udział w sprzedawanych produktach i stąd problemy związane z ich funkcjonowaniem, jak również nowe pojawiające się czynniki wpływające na ich sprzedaż, rentowność i szkodowość są znakomitą obszarą badawczą, tym bardziej, że bogata jest baza danych, gdyż są to częściowo ubezpieczenia obowiązkowe powszechne, o standaryzowanych ogólnych warunkach ubezpieczenia i są masowo sprzedawane. Ponadto trzeba pamiętać, że od 1 maja 2004 r. polski rynek ubezpieczeniowy należy do jednolitego rynku europejskiego i obowiązują dyrektywy UE, które w zakresie ubezpieczeń obowiązkowych OC są dość rozbudowane. W artykule najpierw zostanie przedstawiona krótka charakterystyka całego polskiego rynku ubezpieczeniowego, a następnie – oddzielnie – pewna analiza i tendencje dla ubezpieczeń komunikacyjnych OC i AC. W zakończeniu rozważań zostanie zwrócona uwaga na pewne uwarunkowania oraz perspektywy dalszego rozwoju tych grup ubezpieczeń.

1. Krótka charakterystyka polskiego rynku ubezpieczeniowego

Za początek odrodzenia się rynku ubezpieczeniowego w Polsce można przyjąć datę 28 lipca 1990 r., kiedy uchwalono nową ustawę o działalności ubezpieczeniowej, która uwzględniała już częściowo dyrektywy Unii Europej-

* Uniwersytet Ekonomiczny we Wrocławiu.

skiej dotyczące funkcjonowania jednolitego rynku finansowego i jednocześnie dawała podwaliny prawne pod nowoczesny rynek ubezpieczeniowy działający w warunkach wolnego rynku.

Ważną datą dla rozwoju rynku ubezpieczeniowego w Polsce był dzień 22 maja 2003 r., data uchwalenia pakietu czterech ustaw ubezpieczeniowych, który zaczął obowiązywać polski rynek od 1 stycznia 2004 r. Wszystkie te ustawy uwzględniały rozwiązania proponowane w dyrektywach UE, które miały stworzyć konstrukcję do funkcjonowania jednolitego rynku ubezpieczeniowego w krajach członkowskich UE. Można zatem stwierdzić, że rynek ubezpieczeniowy w Polsce pod względem prawnym jest całkowicie dostosowany do funkcjonowania na wspólnym rynku europejskim. Dużo gorzej to wygląda, gdy weźmiemy pod uwagę wskaźniki ekonomiczne i stopień rozwoju. Zauważyć tu można bezsprzecznie wiele słabości, szczególnie w zakresie wykorzystania usług ubezpieczeniowych i to zarówno przez przedsiębiorstwa, jak również przez gospodarstwa domowe; można mieć też wiele uwag co do jakości tych usług, a także bogactwa oferty proponowanych produktów ubezpieczeniowych. Duże wątpliwości budzi świadomość ubezpieczeniowa Polaków, wiedza kadry kierowniczej i kompetencje pracowników zakładów ubezpieczeń. Cechą charakterystyczną są ciągle niestety niskie kapitały własne, a przede wszystkim duży udział kapitału zagranicznego w kapitałach własnych. Wadą polskiego rynku ubezpieczeniowego jest również zbyt duża koncentracja rynku zarówno w obszarze działu II, jak i działu I. Dominuje na rynku z bardzo dużym udziałem kilka zakładów ubezpieczeń, na czele z grupą PZU, natomiast cała reszta zakładów ma mały udział w rynku – poniżej 5%¹.

Trzeba jednocześnie również zauważyć, że rynek ubezpieczeniowy stanowi integralną część rynku finansowego, stąd bardzo reagował w tym okresie na zmiany, wahania i rozwój rynku kapitałowego, co jest związane z działalnością lokacyjną zakładów ubezpieczeń; reagował również na rozwój rynku bankowego oraz pojawianie się nowych produktów bankowych, proponując uzupełniające produkty ubezpieczeniowe. Od wejścia Polski w struktury UE można zauważyć stopniowy rozwój rynku ubezpieczeniowego, aczkolwiek jest on powolny. Można zaobserwować wiele cech pozytywnych i prawidłowe tendencje w różnych obszarach.

Dobrym miernikiem często stosowanym do oceny rozwoju rynku ubezpieczeniowego w danym kraju i znaczenia sektora ubezpieczeń w gospodarce danego kraju jest procentowy udział składek ubezpieczeniowych w PKB. W tabeli 1 zostanie zaprezentowany ten wskaźnik dla Polski w latach 1992-2009.

¹ J. Handschke, *Polskie doświadczenie w formowaniu i rozwoju rynku ubezpieczeń – wybrane aspekty*, „Wiadomości Ubezpieczeniowe” 2009, nr 3.

Tabela 1

Udział składek ubezpieczeniowych w PKB w Polsce

Rok	Udział w %
1992	1,8
1993	1,9
1994	1,8
1995	1,8
1996	2,1
1997	2,6
1998	2,8
1999	2,9
2000	3,1
2001	3,1
2002	3,01
2006	3,5
2007	3,7
2008	4,7
2009	3,8

Źródło: European Insurance in Figure, Complete 1997 Data, Comite European des Assurances 1998, www.knuife.gov.pl, Polski Rynek Ubezpieczeniowy 2004-2008, GUS, Warszawa 2009.

Dla porównania można podać, że w krajach UE wskaźnik ten wyniósł w 1992 r. 6,3%, a w 1997 r. wyniósł już 7,6%. Natomiast w 2001 r. średnia tego wskaźnika dla krajów UE wynosiła 9,5%, lecz po wejściu do UE państw z Europy Środkowo-Wschodniej średnia tego wskaźnika spadła do 8,3% w 2006 r.

Jeśli weźmiemy pod uwagę brytyjski rynek ubezpieczeniowy – jeden z najbardziej rozwiniętych rynków ubezpieczeniowych na świecie, a w Europie o najwyższym poziomie rozwoju, to w 2006 r. wskaźnik ten wyniósł 16,5%.

Innym wskaźnikiem świadczącym o poziomie świadomości ubezpieczeniowej oraz o zapotrzebowaniu na ochronę ubezpieczeniową i stopniu rozwoju rynku w danym kraju jest udział składki ubezpieczeniowej przypadającej na jednego mieszkańca. Ten wskaźnik w Polsce w 1998 r. wynosił 107,3 USD, a w Wielkiej Brytanii – 2858,9 USD. Natomiast w całej Europie wynosił przeciętnie 613,5 USD. W 2002 r. wskaźnik ten dla Polski wyniósł już 140 USD, a dla Wielkiej Brytanii 3393,8 USD.

Wskaźniki te pokazują jaki jest poziom rozwoju naszego rynku ubezpieczeniowego i jak daleko nam jeszcze do poziomu rozwiniętych rynków ubezpieczeniowych w Europie. Można jednak zaobserwować pewne tendencje

pozytywne, gdyż wskaźniki te sukcesywnie rosną, może w zbyt małym tempie, ale jest to bardzo silnie związane z takimi zjawiskami, jak ogólny rozwój gospodarczy kraju, wzrost PKB czy poprawiający się poziom życia.

W poniższej tabeli jest pokazane jak na przestrzeni ostatnich kilkunastu lat kształtowała się struktura składki w dziale II i jakie pojawiły się zmiany.

Tabela 2

Struktura składki przypisanej brutto według rodzajów działalności w dziale II (w %)

Rok	Składki							
	pozostałe osobowe, gr. 1 + 2	majątkowe, gr. 8 + 9	auto-casco gr. 3	OC komunikacyjne, gr. 10	transportowe, gr. 4 – 7, 11, 12	OC ogólne gr. 13	finansowe gr. 14 – 17	pozostałe gr. 18 – 19
1997	4,9	17,0	33,0	36,2	2,6	2,1	1,8	2,4
1998	4,9	16,7	32,0	37,8	2,2	2,3	2,1	2,1
1999	5,1	16,2	31,5	36,9	1,8	2,7	2,4	3,5
2000	5,3	16,7	30,6	37,3	1,7	2,9	2,7	2,8
2001	5,4	17,7	30,6	36,0	1,8	3,2	3,0	2,3
2002	5,7	18,5	28,7	36,1	1,9	3,4	3,4	2,3
2003	5,7	19,3	30,1	34,0	1,9	3,8	2,5	2,8
2004	5,5	18,4	29,7	33,2	1,7	4,2	4,4	2,9
2005	5,7	17,8	27,8	34,9	1,9	4,5	4,5	2,9
2006	6,0	17,7	25,7	34,7	1,8	5,0	5,5	3,6
2007	6,0	17,7	25,7	34,7	1,8	5,0	5,5	3,6
2008	7,5	16,3	25,5	34,4	1,5	4,8	6,6	3,4
2009	7,3	18,0	23,7	34,6	1,4	5,4	8,2	1,4

Źródło: Raporty PIU, Ubezpieczenia 2006, Ubezpieczenia 2008, Ubezpieczenia 2009.

Jak widać, w dziale II ciągle największy udział w składce zajmują ubezpieczenia komunikacyjne, z zauważalną lekką tendencją zmniejszania się na korzyść innych ubezpieczeń. Na uwagę zasługują ubezpieczenia finansowe, których udział w ciągu 13 lat wzrósł ponad trzykrotnie; jest to związane przede wszystkim z rozwojem działalności kredytowej banków, szczególnie w obszarze udzielania kredytów hipotecznych oraz z upowszechnieniem kart kredytowych i bankomatowych. Ubezpieczenia finansowe są sprzedawane w sieci banków i występują w związku z rozwiniętą współpracą banków oraz ubezpieczycieli. Ponadto przedsiębiorstwa korzystają dość powszechnie z gwarancji ubezpieczeniowych i coraz częściej ubezpieczają się od utraty zysku. Można również zauważyć, że wzrósł ponad dwukrotnie udział ubezpieczenia OC ogólnego. Trzeba pamiętać, że lista ubezpieczeń obowiązkowo-

wych OC wzrosła i aktualnie obowiązuje około 30 ubezpieczeń OC kontraktowych i deliktowych związanych z prowadzeniem różnej działalności gospodarczej oraz uprawianiem niektórych zawodów.

Rynek ubezpieczeniowy od 2004 r. funkcjonuje na jednolitym rynku europejskim i w związku z tym pojawiają się nowe problemy i pewne wyzwania, które są wspólne dla całego rynku europejskiego bądź światowego lub tylko charakterystyczne dla polskiego rynku. Do nowych wyzwań można zaliczyć następujące:

- dalsze zmiany w strukturze sprzedawanych produktów ubezpieczeniowych,
- rozwój ubezpieczeń chorobowych, co jest związane z nowymi formami świadczeń zabezpieczających opiekę chorych,
- dalszy rozwój ubezpieczeń zabezpieczających byt ludzi starszych – nowe strukturyzowane produkty,
- dalsza współpraca z bankami w obszarze dystrybucji produktów i nowych ubezpieczeń finansowych,
- rozwój nowych produktów dla przedsiębiorstw, co jest związane z lepszym zarządzaniem ryzykiem w przedsiębiorstwach,
- udział ubezpieczeń w zarządzaniu katastrofami,
- większa współpraca z innymi instytucjami rynku finansowego, pojawianie się ubezpieczeniowych grup oraz konglomeratów finansowych,
- udoskonalenie metod zarządzania ryzykiem w zakładach ubezpieczeń lub w ubezpieczeniowych grupach kapitałowych, co jest związane z wprowadzeniem nowych zasad nadzoru dla UE (2012 r.).

Tych nowych wyzwań można wymienić więcej, jednak trzeba mieć na uwadze, że możliwości rozwoju rynku ubezpieczeniowego są zawsze silnie powiązane z rozwojem gospodarczym i koniunkturą gospodarczą. Ponadto dane zawarte w tabeli 2 ukazują, że ciągle dominującymi ubezpieczeniami są ubezpieczenia komunikacyjne, zatem ich funkcjonowanie i pewne tendencje, a szczególnie ich rentowność i wyniki finansowe wygenerowane przez te grupy ubezpieczeń będą ciągle wpływać na całokształt działalności ubezpieczeniowej w Polsce.

2. Ubezpieczenia komunikacyjne OC – analiza rozwoju na polskim rynku

Przedmiotem ubezpieczenia komunikacyjnego OC są szkody na mieniu lub osobach spowodowane przez kierujących pojazdem mechanicznym. Jest to ubezpieczenie obowiązkowe powszechne i stąd wszystkie sprzedawane produkty zawierają ogólne warunki ubezpieczenia standardowe określone w Rozporządzeniu Ministra Finansów.

Ubezpieczenia komunikacyjne OC funkcjonują zgodnie z następującymi regulacjami prawnymi:

- Ustawa o działalności ubezpieczeniowej uchwalona 23 maja 2003 r., która weszła w życie 1 stycznia 2004 r.,
- Ustawa o ubezpieczeniach obowiązkowych, Ubezpieczeniowym Funduszu Gwarancyjnym i Polskim Biurze Ubezpieczycieli Komunikacyjnych.

Z rynkiem ubezpieczeń komunikacyjnych OC bardzo silnie są związane dwie instytucje centralne wspomagające ich obsługę, a mianowicie UFG i PBUiK.

Inne regulacje mają charakter międzynarodowy i zobowiązują poszczególne państwa członkowskie UE do wprowadzenia rozwiązań w zakresie ubezpieczenia OC posiadaczy pojazdów mechanicznych oraz są realizowane w postaci tzw. Dyrektyw Komunikacyjnych²:

- I Dyrektywa Komunikacyjna (72/166/EWG) z dnia 24 kwietnia 1972 r. wprowadziła obowiązkowość ubezpieczenia komunikacyjnego OC i zapoczątkowała proces rozwoju i uszczegółowiania prawa na poziomie wspólnotowym;
- II Dyrektywa Komunikacyjna (84/5/EWG) z dnia 30 grudnia 1983 r. wprowadziła obowiązek objęcia ubezpieczeniem OC zarówno szkód na osobie, jak i w mieniu;
- III Dyrektywa Komunikacyjna (90/232/EWG) z dnia 14 maja 1990 r. uściśliła pojęcia i regulacje wprowadzone wcześniejszymi dyrektywami, co wynikało z rozbieżnych interpretacji i praktyki w poszczególnych krajach; wzmocniono tutaj pozycję osób poszkodowanych;
- IV Dyrektywa Komunikacyjna (2000/26/WE) z dnia 16 maja 2000 r. regulowała sytuację osób, które w obcym państwie członkowskim stały się ofiarami wypadków spowodowanych przez pojazdy zarejestrowane w innym państwie niż to, w którym mieszkał poszkodowany, zapewniając poszkodowanym lepszą ochronę ubezpieczeniową;
- V Dyrektywa Komunikacyjna (2005/14/WE) z dnia 11 maja 2005 r. doprecyzowała i rozwinęła przepisy poprzednich dyrektyw; odnosi się przede wszystkim do sfery odszkodowawczej i ma charakter prokonsumencki, porusza zagadnienia związane ze szkodami spowodowanymi przez nieubezpieczonych posiadaczy pojazdów mechanicznych.

Aktualnie w obowiązującym prawie ubezpieczeniowym Polska uwzględniła wszystkie te dyrektywy.

² M. Monkiewicz, Regulacyjna baza jednolitego rynku ubezpieczeń Unii Europejskiej, (w:) Jednolity rynek ubezpieczeń w UE, Procesy rozwoju i integracji, (red.) J. Monkiewicz, Oficyna Wydawnicza Branta, Bydgoszcz-Warszawa 2005.

Ubezpieczenia komunikacyjne OC odgrywają dominującą rolę w większości zakładów ubezpieczeń działu II. Na koniec 2008 r. spośród 35 ubezpieczycieli działu II zezwolenie na prowadzenie działalności ubezpieczeniowej w zakresie ubezpieczeń komunikacyjnych posiadało 27 ubezpieczycieli, a faktycznie działalność tę prowadziło 24 ubezpieczycieli.

Od kilku lat można zaobserwować proces stopniowej dekoncentracji na rynku ubezpieczeń komunikacyjnych OC, co przedstawia poniższa tabela.

Tabela 3

Udział największych ubezpieczycieli w grupie 10 (ubezpieczenia komunikacyjne OC) w Polsce w latach 2002-2009 (w %)

Wyszczególnienie	Lata							
Zakłady	2002	2003	2004	2005	2006	2007	2008	2009/I pół- roczne
1 największy	64,47	60,90	56,13	55,00	47,79	43,36	40,65	39,24
3 największe	80,53	77,11	73,89	72,14	68,06	62,97	58,22	54,00
5 największych	87,45	84,37	80,43	78,99	76,90	73,61	70,04	64,50

Źródło: Rynek ubezpieczeń komunikacyjnych w Polsce i w Unii Europejskiej w latach 2002-2009, Wydawnictwo UFG, Warszawa 2010.

Jak widać, to 5 największych ubezpieczycieli ma przypis składki z tytułu ubezpieczeń komunikacyjnych OC powyżej 50%, jednak jego wielkość w pierwszym półroczu 2009 r. w stosunku do 2002 r. spadła o 22,95%. Największy uczestnik rynku, PZU S.A. w latach 90. był czołowym sprzedawcą polis z zakresu ubezpieczeń komunikacyjnych OC, jednak już jego udział w 2002 r. spadł do 64,47%, a w I półroczu 2009 r. wynosił już tylko 39,24%. PZU S.A. na przestrzeni tych kilku lat od 2002 r. do 2009 r. zanotował spadek udziału w rynku ubezpieczeń komunikacyjnych OC o 25,23%.

Następnie zostanie przedstawiona krótka analiza rynku ubezpieczeń komunikacyjnych OC pod względem wybranych parametrów, takich jak: liczba sprzedanych polis, składka przypisana, wypłacone odszkodowania, udział reasekuracji, szkodowość i rentowność.

Tabela 4

Liczba polis ubezpieczeń komunikacyjnych OC w Polsce w latach 2002-2009 (w mln szt.)

Lata	2002	2003	2004	2005	2006	2007	2008	2009/I pół- roczne
Liczba polis	13,55	12,50	13,32	14,03	14,77	15,76	16,77	16,58

Źródło: Ibid.

W tabeli 4 można zaobserwować ciągły wzrost liczby polis z zakresu ubezpieczenia komunikacyjnego OC, co oczywiście jest związane z rozwojem rynku motoryzacyjnego, jednak dane te nie obrazują w całości zjawiska stosunku liczby polis do liczby pojazdów mechanicznych zarejestrowanych w Polsce, gdyż w praktyce występują tzw. polisy flotowe, w ramach których ubezpieczanych jest kilka, a nawet kilkadziesiąt pojazdów mechanicznych. Zjawisko to jest lepiej przedstawione w następnej tabeli.

Tabela 5

Struktura polis ubezpieczeń komunikacyjnych OC według podmiotów ubezpieczających w Polsce w latach 2002-2009 (w %)

Wyszczególnienie	Lata							
Ubezpieczający	2002	2003	2004	2005	2006	2007	2008	2009/I pół- roczne
osoby fizyczne	92,62	91,80	92,00	91,47	90,31	90,27	89,65	89,33
przedsiębiorstwa	6,90	6,31	6,55	7,10	8,39	8,39	9,12	9,03
pozostałe	0,48	1,89	1,45	1,33	1,30	1,33	1,23	1,64

Źródło: Ibid.

Wskaźnik motoryzacji liczony jako liczba samochodów osobowych na 1 tys. osób w Polsce cały czas rośnie, gdyż w 2002 r. wynosił 289, w 2007 r. wynosił już 383. Jednocześnie warto zwrócić uwagę na fakt, że w Polsce przeważają samochody importowane oraz stare, gdyż w I półroczu 2009 r. zarejestrowano 44,7% samochodów powyżej 10 lat żywotności oraz tylko 11% samochodów importowanych, których wiek nie przekroczył 4 lat.

Innym bardzo istotnym miernikiem funkcjonowania tych ubezpieczeń jest składka przypisana brutto i jej dynamika, co zostanie przedstawione w następnej tabeli dla okresu lat 2002-2008.

Tabela 6

Składka przypisana brutto z tytułu ubezpieczeń komunikacyjnych OC i jej dynamika w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wartość w mld zł	4,27	4,08	4,68	5,40	5,63	6,04	6,81
Dynamika w %		95,55	114,71	115,38	104,26	107,28	112,75

Źródło: Na podstawie: Ibid.

Dynamika zbioru składki miała przeważnie tendencję wzrostową. Następna tabela przedstawia jak kształtowała się średnia wysokość składki z ubezpieczeń komunikacyjnych OC w latach 2002-2008.

Tabela 7

Średnia wysokość składki w ubezpieczeniach komunikacyjnych OC i jej dynamika w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wartość w zł	322,65	326,51	351,52	384,75	381,60	383,02	406,45
Dynamika w %		101,20	107,66	109,45	99,18	100,37	106,12

Źródło: Na podstawie: Ibid.

Jak widać, składka cały czas rosła z wyjątkiem 2006 r. Wzrost średniej składki brutto występował równocześnie ze wzrostem liczby polis oraz z pewnymi procesami zachodzącymi w polskiej gospodarce, z wejściem Polski na jednolity rynek UE oraz ze zmianami na polskim rynku motoryzacyjnym m.in. takimi jak wzrost szkodowości i wprowadzeniem tzw. „podatku Religi”, który przestał obowiązywać z dniem 1 stycznia 2009 r.

Następna tabela ukazuje, jak kształtował się poziom odszkodowań i świadczeń brutto z ubezpieczeń komunikacyjnych OC wraz z dynamiką tego zjawiska w analogicznym okresie 2002-2008.

Tabela 8

Odszkodowania i świadczenia brutto z ubezpieczeń komunikacyjnych OC oraz ich dynamika w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wartość w mld zł	2,80	2,85	3,13	3,21	3,43	3,87	4,38
Dynamika w %		101,79	109,82	102,56	106,85	112,83	113,18

Źródło: Na podstawie: Ibid.

Można również zaobserwować wzrost częstości wypłat z ubezpieczeń komunikacyjnych OC w rozważanych latach, który w 2002 r. wynosił 4,46%, a w 2008 r. kształtował się na poziomie 5,14%, co powinno skutkować podwyższeniem stopy składki.

Poziom reasekuracji jest mierzony wskaźnikiem zatrzymania składki i wskaźnikiem zatrzymania odszkodowań na udziale własnym. Tutaj dla ubezpieczeń komunikacyjnych OC można zauważyć dla obu tych wskaźników wyraźną tendencję wzrostową, co świadczy o zmniejszonym udziale reasekuracji biernej w tej grupie ubezpieczeń, tym bardziej, że warto wspomnieć lata 90., kiedy to wysokość wskaźnika, a szczególnie zatrzymania składki, była bardzo niska – oscylowała w granicach 60-70%. Zjawisko to obrazuje tabela 9.

Tabela 9

Wskaźnik zatrzymania składki i odszkodowań dla ubezpieczeń komunikacyjnych OC
w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wskaźnik zatrzymania składki w (%)	73,09	85,64	84,11	85,12	90,06	92,05	92,09
Wskaźnik zatrzymania odszkodowań w (%)	82,39	80,74	82,02	83,62	83,58	87,91	92,56

Źródło: Ibid.

Warto zauważyć, że te dwa wskaźniki kształtowały się przeważnie na zbliżonym poziomie, co świadczy o dobrej polityce reasekuracyjnej dla tej grupy ubezpieczeń.

Na zakończenie analizy statystycznej ubezpieczeń komunikacyjnych OC funkcjonujących na polskim rynku pokażemy jak kształtował się w analogicznym okresie badawczym wskaźnik rentowności technicznej oraz wskaźnik szkodowości na udziale własnym.

Tabela 10

Wskaźniki rentowności i szkodowości na udziale własnym z ubezpieczeń komunikacyjnych OC
w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wskaźnik rentowności z działalności technicznej w (%)	0,59	1,95	-4,59	1,59	4,07	-2,73	-11,31
Wskaźnik szkodowości na udziale własnym w (%)	75,20	77,70	78,68	72,64	72,25	79,73	78,99

Źródło: Ibid.

Te dwa wskaźniki obrazują niepokojące zjawisko, które pojawiło się w ostatnich latach w tej grupie ubezpieczeń. Jak widać, ubezpieczenia komunikacyjne OC w ostatnich latach przynoszą straty z działalności technicznej, są zupełnie nierentowne i jednocześnie charakteryzują się tendencją wzrostową w zakresie szkodowości. Można to tłumaczyć takimi zjawiskami, jak wzrost przestępczości ubezpieczeniowej lub za niskie składki. Z pewnością z problemem tym w najbliższym czasie będą musiały sobie poradzić zakłady ubezpieczeń z działu II, które mają duży udział tej grupy ubezpieczeń w swoim portfelu sprzedawanych produktów.

3. Analiza i tendencje funkcjonowania ubezpieczenia komunikacyjnego AC na polskim rynku

Ubezpieczenie komunikacyjne autocasco (grupa 3) jest dobrowolne i ma na celu bezpośrednią ochronę majątku ubezpieczonego, z chwilą zaistnienia zdarzenia ubezpieczeniowego odszkodowanie wypłacane z tego ubezpieczenia ma za zadanie pokryć szkodę będącą następstwem takich zdarzeń, jak zniszczenie, uszkodzenie i kradzież pojazdu. Przedmiotem ubezpieczenia są pojazdy wraz z wyposażeniem, z tym że zdecydowana większość ogólnych warunków ogranicza ochronę ubezpieczeniową do wyposażenia standardowego, jednak przewiduje możliwość rozszerzenia ochrony za dodatkową składkę.

Z uwagi na fakt, że jest to ubezpieczenie dobrowolne, to ogólne warunki ubezpieczenia określa ubezpieczyciel, dlatego mogą one różnić się u poszczególnych ubezpieczycieli. Ponadto, ponieważ nie obowiązuje przymus ubezpieczenia, to nie wszystkie pojazdy zarejestrowane posiadają to ubezpieczenie, jednak w ostatnich latach coraz więcej samochodów kupowanych jest na kredyt, stąd banki wymagają od w ten sposób sprzedanych samochodów ubezpieczenia AC.

Dokonyamy przeglądu analizy statystycznej funkcjonowania ubezpieczenia AC, podobnie jak to zostało zrobione wyżej dla ubezpieczenia OC i porównamy wyniki. Istotne znaczenie w kształtowaniu szkodowości w ubezpieczeniach komunikacyjnych AC ma zjawisko kradzieży samochodów. Warto zauważyć, że zjawisko to znacznie uległo poprawie, gdyż z roku na rok notuje się mniejszą liczbę skradzionych samochodów. W roku 2002 skradziono 53 674 samochodów, a w 2008 r. już tylko 17 669 samochodów.

Najpierw przyjrzyjmy się koncentracji rynku w tej grupie ubezpieczeń.

Tabela 11

Udział największych ubezpieczycieli w grupie 10 (ubezpieczenia komunikacyjne AC) w Polsce w latach 2002-2009 (w %)

Wyszczególnienie	Lata							
	2002	2003	2004	2005	2006	2007	2008	2009/I pół- roczne
1 największy	64,10	62,60	60,75	60,82	59,30	55,45	50,54	46,22
3 największe	83,84	82,27	79,09	79,19	77,74	74,46	71,21	68,71
5 największych	89,90	89,34	87,49	88,21	87,75	86,62	83,73	81,87

Źródło: Ibid.

W tej grupie ubezpieczeń proces dekoncentracji rynku przebiega bardzo powoli i w zasadzie nie można mówić o rozproszonym rynku, gdyż pierwsza 5 ma cały czas udział w składce przypisanej powyżej 80%. Natomiast lider rynku PZU S.A. sukcesywnie zmniejsza udział w rynku z roku na rok, jednak w tej grupie ubezpieczeń jest ciągle podstawowym ubezpieczycielem z udziałem powyżej 40%.

Tabela 12

Liczba polis ubezpieczeń komunikacyjnych AC w Polsce w latach 2002-2009
(w mln szt.)

Lata	2002	2003	2004	2005	2006	2007	2008	2009/I pół- roczne
Liczba polis	3,94	3,80	3,86	3,93	3,93	4,25	4,58	4,61

Źródło: Ibid.

Liczba sprzedanych polis cały czas rośnie, co jest oczywiste w związku z rozwojem motoryzacji i rosnącą liczbą zarejestrowanych samochodów, jednak w porównaniu z polisami ubezpieczenia OC jest to liczba około 3 razy mniejsza, co oznaczałoby, że co trzeci samochód nie posiada ubezpieczenia AC.

Tabela 13

Struktura polis ubezpieczeń komunikacyjnych AC według podmiotów ubezpieczających w Polsce w latach 2002-2009 (w %)

Wyszczególnienie	Lata							
Ubezpieczający	2002	2003	2004	2005	2006	2007	2008	2009/I pół- roczne
osoby fizyczne	83,08	82,20	80,77	78,86	74,15	73,64	71,88	72,48
przedsiębiorstwa	15,94	15,07	16,06	18,19	22,70	23,10	25,35	24,90
pozostałe	0,98	2,73	3,17	2,95	3,15	3,26	2,77	2,61

Źródło: Ibid.

Można zauważyć zjawisko rosnącej tendencji zwiększania się udziału przedsiębiorstw w liczbie zawartych umów ubezpieczeń komunikacyjnych AC. Porównując te dane z analogicznymi do ubezpieczeń komunikacyjnych OC, można zauważyć, że struktura sprzedanych polis dla osób fizycznych i przedsiębiorstw jest korzystniejsza dla przedsiębiorstw w przypadku ubezpieczenia AC.

Tabela 14

Składka przypisana brutto z tytułu ubezpieczeń komunikacyjnych AC i jej dynamika w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wartość w mld zł	3,79	4,04	4,43	4,36	4,22	4,74	5,21
Dynamika w %		106,60	109,65	98,42	96,79	112,32	109,92

Źródło: Na podstawie: Ibid.

Dynamika wzrostu składki w ubezpieczeniach komunikacyjnych AC charakteryzowała się zmiennością, jednak miała tendencję wzrostową. W 2008 r. składka przypisana brutto była o 37% wyższa niż w 2002 r. Wzrost składki można wytłumaczyć takimi zjawiskami, że Polacy kupują nowsze i droższe samochody oraz zwiększyła się sprzedaż ratalna, która wymusza zakup ubezpieczenia AC.

Tabela 15

Średnia wysokość składki w ubezpieczeniach komunikacyjnych AC i jej dynamika w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wartość w zł	989,48	1 062,55	1 147,74	1 109,16	1 075,17	1 114,27	1 136,79
Dynamika w %		107,38	108,02	96,64	96,94	103,64	102,02

Źródło: Na podstawie: Ibid.

Średnia wysokość składki w ubezpieczeniach AC jest znacznie wyższa od ubezpieczenia OC, co wynika z samej konstrukcji tego ubezpieczenia i znacznie wyższej częstotliwości występowania wypadków ubezpieczeniowych. Dynamika średniej wysokości składki była bardzo podobna do dynamiki całej składki przypisanej, co jest oczywiste, można tylko zanotować, że jej wzrost w 2008 r. w stosunku do 2002 r. był mniejszy niż powyżej, gdyż wyniósł 14,89%.

Zostanie przedstawione jak kształtowały się odszkodowania i świadczenia wypłacone z tytułu ubezpieczenia komunikacyjnego AC w latach 2002-2008 oraz jaka była dynamika tego zjawiska.

Tabela 16

Odszkodowania i świadczenia brutto z ubezpieczeń komunikacyjnych AC oraz ich dynamika w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wartość w mld zł	2,97	2,98	3,14	3,10	2,90	2,79	3,11
Dynamika w %		100,34	105,37	98,73	93,55	96,21	111,47

Źródło: Na podstawie: Ibid.

Dynamika wypłat odszkodowań i świadczeń charakteryzowała się wahaniami, z tym że początkowo w badanym okresie zanotowano lekki wzrost, następnie spadek, a w ostatnim roku (2008) w stosunku do poprzedniego zanotowano wzrost o 11,47%.

Inny bardzo interesujący wskaźnik to częstość wypłat z ubezpieczeń komunikacyjnych AC, który w porównaniu ze wskaźnikiem dla ubezpieczeń komunikacyjnych OC jest dużo wyższy, jednak w ostatnich latach miał tendencję spadkową. W 2002 r. wynosił dla ubezpieczeń AC 16,36%, a w 2008 r. analogiczny wskaźnik wyniósł 14,11%. Tak wysoki wskaźnik wpływa na znacznie wyższą stopę składki ubezpieczeniowej dla ubezpieczeń komunikacyjnych AC.

Następnie zostanie przedstawione jaki był udział w reasekuracji biernej ubezpieczycieli w zakresie prowadzenia ubezpieczeń komunikacyjnych AC i czy prawidłowa była polityka reasekuracyjna.

Tabela 17

Wskaźnik zatrzymania składki i odszkodowań dla ubezpieczeń komunikacyjnych AC w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wskaźnik zatrzymania składki w (%)	92,60	94,45	94,09	95,35	95,26	96,59	98,13
Wskaźnik zatrzymania odszkodowań w (%)	87,54	93,41	93,39	94,78	96,14	97,29	98,30

Źródło: Ibid.

Dane w tabeli pokazują, że polityka reasekuracyjna polskich ubezpieczycieli w stosunku do ubezpieczeń komunikacyjnych AC jest poprawna oraz udział reasekuracji biernej jest na bezpiecznym poziomie i jednocześnie jest zdecydowanie niższy dla tej grupy ubezpieczeń niż dla ubezpieczeń komunikacyjnych OC.

Na zakończenie przeprowadzonej analizy i powyższych rozważań przyjrzyjmy się, jak kształtowały się w badanym okresie wskaźniki rentowności i szkodowości na udziale własnym dla ubezpieczeń komunikacyjnych AC i porównajmy je z ubezpieczeniami OC.

Tabela 18

Wskaźniki rentowności i szkodowości na udziale własnym z ubezpieczeń komunikacyjnych AC w latach 2002-2008

Lata	2002	2003	2004	2005	2006	2007	2008
Wskaźnik rentowności z działalności technicznej (w %)	-7,04	-9,01	-1,99	1,55	5,12	7,04	8,90
Wskaźnik szkodowości na udziale własnym (w %)	81,00	77,96	71,72	69,75	67,10	64,87	63,20

Źródło: Ibid.

Jak widać z powyższej tabeli, ubezpieczenia komunikacyjne AC poczynawszy od 2005 r. są znacznie bardziej rentowne niż ubezpieczenia OC, i tak w 2008 r. wygenerowały wynik techniczny dodatni na poziomie 439 mln zł. Poczynawszy od 2005 r. ta grupa ubezpieczeń (AC komunikacyjne) jest znacznie bardziej opłacalna. Natomiast w przypadku ubezpieczenia OC pojawiały się duże wahania wyniku technicznego, a w ostatnich latach strata techniczna znacznie się pogłębiła. Na uwagę zasługuje również fakt, że znacznie poprawił się wskaźnik szkodowości w ubezpieczeniach komunikacyjnych AC, który ma tendencję spadkową.

4. Uwarunkowania i perspektywy dalszego funkcjonowania ubezpieczeń komunikacyjnych w Polsce

Analizując powyższe dane, można wyciągnąć następujące wnioski co do funkcjonowania ubezpieczeń komunikacyjnych na polskim rynku:

- duże zróżnicowanie między ubezpieczeniami OC a ubezpieczeniami AC; OC prawie nierentowne z dużą szkodowością, natomiast ubezpieczenia AC rentowne i ze zmniejszającą się szkodowością,
- znacznie wyższa częstotliwość występowania szkód w ubezpieczeniach AC niż w ubezpieczeniach OC,
- znacznie wyższa przeciętna składka w ubezpieczeniach AC niż w ubezpieczeniach OC,

- powyższe problemy nie są tak widoczne, gdyż zakłady ubezpieczeń rozpraszają ryzyko funkcjonowania tych ubezpieczeń w praktyce, sprzedając ubezpieczenia pakietowe,
- około 30% zarejestrowanych samochodów posiada ubezpieczenie AC,
- duży udział procentowy wśród zarejestrowanych samochodów stanowią samochody wieloletnie, co wpływa na ich niską wartość oraz niższą składkę ubezpieczenia AC,
- zmniejszający się udział reasekuracji w tych grupach ubezpieczeń.

Ubezpieczenia komunikacyjne są w Polsce ciągle podstawowymi produktami sprzedawanymi przez zakłady ubezpieczeń, gdyż ich udział w dziale II jest bardzo wysoki w porównaniu do innych krajów rozwiniętych gospodarczo w Europie. Przykładowo zostaną zaprezentowane dane dla 2007 r., w którym udział ubezpieczeń komunikacyjnych w Polsce w dziale II wyniósł 59,8%, natomiast w Niemczech tylko 23,7%, w Wielkiej Brytanii 26%, w Czechach 49,1%, na Litwie 65,4%, a w Rumunii 79,8%. Z drugiej strony, jeśli popatrzymy na wspomniany wcześniej wskaźnik motoryzacji (liczba samochodów osobowych na 1000 mieszkańców), to w Polsce jest na niższym poziomie niż w krajach UE, gdyż w tymże 2007 r. wyniósł 383 w Polsce, we wszystkich krajach UE – 464, w UE-15 – 500, w Niemczech – 501, a w Wielkiej Brytanii – 476. W związku z tymi danymi myślę, że tak wysoki udział ubezpieczeń komunikacyjnych na polskim rynku ubezpieczeniowym jest ciągle związany z małym zainteresowaniem innymi produktami ubezpieczeniowymi, na które popyt jest ciągle niski, co wynika ze słabego rozwoju rynku, a przede wszystkim z niższym poziomem rozwoju gospodarczego, gdyż ubezpieczenia są bardzo wrażliwe na koniunkturę gospodarczą.

Obserwując dane statystyczne, można zauważyć, że tendencje są prawidłowe i wszystkie te procesy, które zbliżają nasz rynek ubezpieczeniowy do rynków w krajach wysoko rozwiniętych, zachodzą, ale bardzo powoli.

TRANSPORT INSURANCES AS BASIC PRODUCTS SOLD AT POLISH INSURANCE MARKET

Summary

The article considers the problem of insurance operation in Poland. It was noted that their participation in the Polish insurance market is still high compared to other EU countries with a high level of economic development. Conducted a brief statistical analysis separately for the insurance and liability insurance AC and compared these two groups, especially their profitability and loss ratio, and some trends associated with such processes as the impact of insurance premiums, the flow of claims and the use of reinsurance for these classes. The article ends with conclusions and observations about Polish conditions and future development.

Włodzimierz Szkutnik

WPŁYW RYZYKA STOPY BAZOWEJ I MORALNEGO HAZARDU NA CENĘ I WYPŁATĘ OBLIGACJI CAT

Wprowadzenie

W artykule tematycznie nawiązuje się do wyceny obligacji CAT¹, z uwzględnieniem moralnego hazardu i ryzyka bazowego instrumentów ubezpieczeniowych, i bez uwzględnienia tych ryzyk². Wycena takich obligacji ma na celu zarządzanie ryzykiem ubezpieczeniowym o charakterze katastrofalnym. Podstawą prowadzonych rozważań będą instrumenty finansowe, których zastosowanie w procesie zarządzania firmą ubezpieczeniową ma na celu sekurytyzację ryzyka katastrofalnego. Znaczący to, że ich celem jest finansowanie skutków zdarzeń katastrofalnych. Narzędziami finansowania będą obligacje CAT. Mają one charakter zbliżony do kapitału warunkowego CC (ang. *contingent claim*), ale w realnych warunkach są instrumentami rynku finansowego – obligacjami katastrofalnymi.

Celem prowadzonych tu rozważań jest sformułowanie modelu warunkowych roszczeń dla wyceny obligacji emitowanych pośrednio przez ubezpieczyciela, poprzez spółkę SPV, gdy istnieje ryzyko odmowy finansowego uregulowania roszczeń zainicjowanych zdarzeniem katastrofalnym. W realiach rynkowych wycena takich obligacji znajdowała uzasadnienie w ostatnich dziesięcioleciach w wyniku zwiększonej częstości występowania katastrofalnych zdarzeń. Przewidywanie skutków ewentualnych katastrof o dużych zagregowanych stratach, wyrażona w pośredni sposób w specyficznie skonstruowanych wskaźnikach (indeksach) strat z uwzględnieniem ich w warunkach emitowanych obligacji, jest antycypacją ich ceny. Skracając istotę problemu, można stwierdzić, że ich właściwa wycena ma wpływ na straty ubezpieczyciela, gdy zostanie zrealizowane zdarzenie katastrofalne. Zatem ich właściwa wycena jest elementem zarządzania działalnością ubezpieczeniową. Wtedy bowiem następuje całkowite lub częściowe umorzenie nominalnej wartości obligacji CAT, co ma wpływ na finanse firmy ubezpieczeniowej.

¹ Por. Jin-Ping Lee, Min-Teh Yu, *Pricing Default-Risky CAT Bonds with Moral Hazard and Basis Risk*, „Journal of Risk and Insurance” 2002, Vol. 69.

² W. Szkutnik, *Anticipation of the Price of CAT Bond in the Management of the Process of Securitization of Catastrophe Insurance*, (w:) *Ekonometria 28. Forecasting*, (red.) Paweł Dittmann, Research Papers of Wrocław of Economics, No. 91, UE, Wrocław 2010, s. 151-165.

Specyficzny projekt obligacji CAT ma kontekst historyczny. W retrospektywie konstrukcji takich projektów obserwowana była nieustanna ewolucja projektów CAT, a zatem związana z tym ich wycena. Choć można wyceniać obligacje CAT w warunkach ryzyka odmowy wypłaty odszkodowań, możliwe jest także uwzględnienie warunków, gdy nie ma takiego ryzyka³. Będziemy także zakładać, że istnieje ryzyko moralne (hazard moralny), z czym łączy się nieadekwatne do strat zjawisko wyceny roszczeń, prowadzące do strat ubezpieczyciela (najczęściej) albo posiadaczy obligacji i polis ubezpieczeniowych. Należy bowiem pamiętać, że częściowo ryzyko katastrofalne jest ubezpieczane w tradycyjny sposób przez podmioty narażone na takie ryzyko i w procedurze tej biorą udział także towarzystwa ubezpieczeniowe, typowo reasekuracyjne. Uwzględnione będzie także ryzyko bazowe, będące symptomem wpływu uwarunkowań rynku kapitałowego i ryzyka rynkowego.

W modelu dynamiki wartości aktywów firmy ubezpieczeniowej istotne założenia będą przyjęte dla ryzyka stopy procentowej i ryzyka kredytowego oddziałującego w określony sposób na wszystkie inne formy ryzyka. Bardziej szczegółowo wyjaśnione to zostanie przy specyfikacji założeń modelu wartości aktywów.

Zwraca uwagę fakt, że zaprezentowane w artykule przykłady empiryczne⁴ mogą przystawać do polskich realiów przejawiania się ryzyka katastrofalnego, chociaż sam rynek kapitałowy nie wypracował jeszcze odpowiednich projektów instrumentów znajdujących zastosowanie w sekurytyzacji takiego ryzyka. W literaturze krajowej podjęto jedynie próbę wyceny obligacji katastrofalnej w ujęciu tzw. dwuczynnikowej funkcji użyteczności inwestora, której remitentem miały być władze komunalne.

1. Strukturalny aspekt obligacji CAT

W polskiej literaturze wspomniane obligacje katastrofalne były rozpatrywane jako instrumenty dłużne obciążone ryzykiem niewypłacalności remitenta. W polskich warunkach projekt takich instrumentów był adresowany do władz samorządowych dla zabezpieczenia regionu przed skutkami powodzi, suszy pożaru lasu itp. Struktura takich obligacji, a stąd też ich skuteczność polega na tym, że w przypadku wystąpienia zjawiska katastrofalnego stymulującego wyprzedzająco emisję takich papierów dłużnych, które generuje straty finansowe i roszczenia, remitent tych papierów nie ma obowiązku zwrotu wartości nominalnej obligacji i kontynuacji wypłaty odsetek do momentu wygaśnięcia terminu ich ważności. Natomiast gdy nie zaistnieje katastrofa, właściciel obligacji (inwestor) jest beneficjentem takiego sprzyjającego układu.

³ Ryzyko zobowiązań i działań w inwestycjach gospodarczych oraz ubezpieczeniowych. Aspekt modelowania i oceny sekurytyzacji ryzyka, (red.) W. Szkutnik, AE, Katowice 2009.

⁴ Ibid.

Istotne cechy obligacji katastroficznych z rynku USA określających ich strukturę to:

- 1) istnienie uruchomienia umorzenia długu; wynika to z warunków emisyjnych obligacji CAT, których ogólny projekt pozwala na wypłatę odsetek i/lub zwrot umorzenia kapitału, przy czym zakres umorzenia może być całkowity lub częściowy, lub wyskalowany do wielkości straty;
- 2) umorzenie długu może być uruchomione przez rzeczywiste straty ubezpieczyciela (lub reasekuratora) lub z obserwacji indeksu start ubezpieczyciela w trakcie ustalonego okresu;
- 3) sekurytyzujące obligacje CAT stwarzają ubezpieczycielowi (reasekuratorowi) szansę uniknięcia ryzyka kredytowego (sposób pozyskania kapitału);
- 4) obligacje CAT gwarantują ubezpieczycielowi zabezpieczenie, poprzez umorzenie istniejącego długu (ubezpieczyciela); zatem wartość zabezpieczenia jest niezależna od aktywów posiadaczy obligacji i remitent obligacji nie ponosi ryzyka „niezrealizowania się zabezpieczenia”;
- 5) ocena wartości obligacji CAT z wykorzystaniem czynników ją determinujących z perspektywy posiadacza obligacji (inwestora), która jest uzasadniona:
 - a) ryzykiem odmowy wypłaty,
 - b) potencjalnej tendencji do ujawnienia się lub „zachowania się” moralnego hazardu,
 - c) ryzykiem bazowym firmy emitującej obligacje;
- 6) hazardem moralnym, który może zwiększyć wypłaty roszczeń kosztem zmniejszenia wartości kapitału inwestorów i może wpłynąć na wycenę obligacji;
- 7) ryzykiem bazowym obligacji CAT odnoszącym się do luki pomiędzy rzeczywistą stratą ubezpieczyciela a łącznym indeksem strat, która powstrzymuje ubezpieczyciela od zawierania kompleksowego zabezpieczenia ryzyka.

2. Tendencje w wycenie projektów obligacji CAT z uwzględnieniem ryzyka moralnego i ryzyka bazowego

Wpływ hazardu moralnego i ryzyka bazowego był badany początkowo nie z perspektywy wyceny obligacji CAT, lecz samego wpływu tych rodzajów ryzyka na proces sekurytyzacji jako formy finansowania rozmaitych przedsięwzięć gospodarczych, niepewnych, a obciążonych znaczącymi stratami w sytuacji niepowodzenia tych przedsięwzięć. Tworzone były też potencjalnie nowe instrumenty z ograniczeniami, związane z sekurytyzacją skutków zjawisk

katastrofalnych⁵. Rozważania z perspektywy korygowania określonych działań i zastosowań z perspektywy zarządzania ryzykiem prowadzili Bouzouita i Young⁶.

Samą wycenę transakcji terminowych CAT *futures* i opcji *spreads* CAT przy założeniu deterministycznej stopy procentowej i specyficznego procesu strat w szacowaniu roszczeń majątkowych (ang. *property claims services*) PCS oceniali m.in. Cox i Schwebach, Cummins i Geman i Chang i Yu⁷. Wyceniana też była jednoroczna zerokuponowa obligacja CAT⁸, która była porównywana z ceną obligacji CAT oszacowaną przez rozkład hipotetycznej straty katastrofalnej. Zajdenweber⁹ poszedł w kierunku Litzenbergera, Beaglehole i Reynoldsa, ale zmienił rozkład straty katastrofalnej na stabilny rozkład Levy'ego. Louberge, Kellezi i Gilli oszacowali numerycznie cenę obligacji CAT przy założeniach, że strata katastrofalna jest czystym procesem Poissona, indywidualne straty mają niezależne, identyczne rozkłady lognormalne, i że model stopy procentowej jest procesem losowym dwumianowym¹⁰.

W żadnym z powyższych opracowań modeli wyceny nie udało się wprowadzić powszechnie akceptowanego stochastycznego procesu stopy procentowej i procesu straty katastroficznej oraz uwzględnienia ryzyka odmowy zapłaty w obligacjach CAT.

Dopiero w pracy Jin-Ping Lee i Min-Teh Yu¹¹ został sformułowany model warunkowych roszczeń dla wyceny obligacji emitowanych przez ubezpieczyciela przy istnieniu ryzyka odmowy finansowego uregulowania roszczeń związanych z katastrofą, w którym stopy procentowe mają charakter stochastyczny. Ponadto dopuszcza się bardziej ogólne procesy strat oraz zezwala na praktyczne uwzględnienie moralnego hazardu, ryzyka stopy bazowej i ryzyka odmowy zapłaty. Szacowania cen obligacji CAT dokonuje się zarówno przy założeniu braku ryzyka odmowy wypłaty odszkodowań, jak i w warunkach istnienia ryzyka odmowy takiej wypłaty. Wyniki tych szacunków wskazują, że zarówno moralny hazard, jak i ryzyko stopy bazowej

⁵ M.S. Canter, J.B. Cole, R.I. Sandor, *Insurance Derivatives: A New Asset Class for the Capital Markets and a New Hedging Tool for the Insurance Industry*, „Journal of Applied Corporate Finance” 1997, Vol. 10, s. 69-83.

⁶ R. Bouzouita, A.J. Young, *Catastrophe Insurance Option*, „Journal of Insurance Regulation” 1998, No. 16(3), s. 313-326.

⁷ S.H. Cox, R.G. Schwebach, *Insurance Futures and Hedging Insurance Price Risk*, „Journal of Risk and Insurance” 1992, No. 59, s. 628-644; J.D. Cummins, H. Geman, *Pricing Catastrophe Futures and Call Spreads: An Arbitrage Approach*, „Journal of Fixed Income” 1995, No. 4, s. 46-57; C.J. Chang, M.T. Yu, *Pricing Catastrophe Insurance Futures Call Spreads: A Randomized Operational Time Approach*, „Journal of Risk and Insurance” 1996, No. 63(4), s. 599-617.

⁸ R.H. Litzenberger, D.R. Beaglehole, C.E. Reynolds, *Assessing Catastrophe Reinsurance-Linked Securities as a New Asset Class*, „Journal of Portfolio Management” 1996, No. 23(3), s. 76-86.

⁹ D. Zajdenweber, *The Valuation of Catastrophe-Reinsurance-Linked Securities*, paper presented at American Risk and Insurance Association Meeting, 1998.

¹⁰ H. Loubergé, E. Kellezi, M. Gilli, *Using Catastrophe-Linked Securities to Diversify Insurance Risk: A Financial Analysis of CAT Bonds*, „Journal of Insurance Issues” 1999, No. 22(2), s. 125-146.

¹¹ Jin-Ping Lee, Min-Teh Yu, op. cit.

obniżają w znacznym stopniu wartość obligacji. Nie powinny być one zatem ignorowane w wycenie obligacji CAT. Wskazuje się także na relacje między cenami obligacji a skalą roszczeń spowodowanych zdarzeniem katastrofalnym, zmiennością straty, poziomem uruchomienia (ang. *level trigger*), zależnym od indeksu strat, pozycją kapitałową firmy emitującej obligacje, strukturą zadłużenia i niepewnością stopy procentowej. Opracowanie tego tematu jest ważne ze względów praktycznych. Wynika to stąd, że przy przyjętych założeniach w proponowanych modelach wartości aktywów, stopy procentowej, modelu strat, wycenione zabezpieczające obligacje CAT umożliwiają emitentowi uniknięcie ryzyka kredytowego, jakie może się pojawić przy tradycyjnej reasekuracji lub przy stosowaniu zabezpieczeń w postaci opcji katastrofalnych.

Nawiązując jeszcze do moralnego hazardu, należy pamiętać, że jest on inicjowany przez samego ubezpieczyciela¹². Związane jest to z kosztem szacowania strat przez ubezpieczyciela. Niekiedy koszty te przewyższają zyski emitenta (spółka SPV) obligacji (ubezpieczyciela) wynikające z wartości długu (emitenta obligacji), jaki pojawia się w momencie zaistnienia zdarzenia katastrofalnego. Wynika to ze specyfiki obligacji CAT. Znaczący to, że ubezpieczyciel przejawia inicjatywę do płacenia bardziej hojnych roszczeń, kiedy wartość straty jest bliska wartości uruchomienia umorzenia długu ustalonego w postanowieniach emisji obligacji. Doherty wskazał, że moralny hazard wynika z mniejszej dbałości o autentyczną ocenę szacowanych i kontrolowanych strat dokonywanych przez ubezpieczyciela emitującego obligacje CAT, ponieważ zwiększenie tej dbałości wpływa na zwiększenie wartości długu, jaki musi być spłacony¹³. Natomiast Bantwal i Kunreuther także zauważyli tendencję ubezpieczycieli do zawierania dodatkowych polis w obszarze narażonym na katastrofy, poświęcając mniej czasu i środków finansowych na szacowanie i kontrolę finansową strat po katastrofie¹⁴. Wpływ moralnego hazardu może zwiększyć wypłaty roszczeń kosztem zmniejszenia wartości kapitału posiadaczy polis i może wpłynąć na cenę obligacji.

Innym ważnym aspektem, jaki musi być uwzględniony w wycenie obligacji CAT jest ryzyko bazowe. Jak wiadomo¹⁵, ryzyko bazowe obligacji CAT odnosi się do luki pomiędzy rzeczywistą stratą ubezpieczyciela a złożonym indeksem strat, która powstrzymuje ubezpieczyciela od zawierania komplekso-

¹² B. Lawędzia, W. Szkutnik, *Aspekt sekurytyzacji i hazardu moralnego w prognozowaniu działalności lokacyjnej firmy ubezpieczeniowej*, (w:) *Prognozowanie w zarządzaniu firmą*, (red.) P. Dittmann, J. Krupowicz, Nr 1112, AE, Wrocław 2006; Eidem, *Ryzyko realizacji wysokich roszczeń aspekcie modelowej równowagi hazardu moralnego i ryzyka bazowego*, (w:) *Metody matematyczne, ekonometryczne i informatyczne w finansach i ubezpieczeniach*, cz. 1, (red.) P. Chrzan, AE, Katowice 2006.

¹³ N. Doherty, *Financial Innovation in the Management of Catastrophe Risk*, „Journal of Applied Corporate Finance” 1997, Vol. 10(3).

¹⁴ V.J. Bantwal, H.C. Kunreuther, *A CAT Bond Premium Puzzle?*, Working paper, Financial Institutions Center, The Wharton School, University of Pennsylvania 1999.

¹⁵ *Zastosowanie metod ekonometryczno-statystycznych w zarządzaniu finansami w zakładach ubezpieczeń*, (red.) Ronka-Chmielowiec, AE, Wrocław 2004.

wego zabezpieczenia ryzyka. Ryzyko bazowe może spowodować, że ubezpieczyciele zwlekają ze spłatą ich długu w przypadku wysokiej indywidualnej straty, ale niskiego wskaźnika straty, a zatem wpływa to na cenę obligacji. Istnieje jednak równowaga pomiędzy ryzykiem bazowym i moralnym hazardem. Jeśli używa się niezależnie obliczonego indeksu dla określenia płatności z obligacji CAT, wówczas możliwość oszukania posiadaczy obligacji przez ubezpieczyciela zmniejsza się lub jest całkowicie wyeliminowana. Jest to równoważne z ujawnieniem zachowania się w zmniejszonym zakresie hazardu moralnego lub też z jego zupełną eliminacją. Pojawia się jednak wtedy ryzyko bazowe.

3. Koncepcja stochastycznej struktury wyceny obligacji CAT

Struktura wyceny obligacji CAT¹⁶ wychodzi naprzeciw nieuwzględnianym z reguły założeniom w modelach wyceny obligacji CAT. Model szacowania wyceny obligacji CAT jest zaprezentowany względem praktycznych założeń dotyczących:

- ryzyka odmowy zapłaty,
- moralnego hazardu,
- ryzyka bazowego.

W modelu tym niezbędne jest określenie:

- dynamiki wartości aktywów A_t ,
- dynamiki stopy procentowej r_t ,
- dynamiki łącznej straty $C_{i,t}$ dla firmy i emitującej obligacje oraz odpowiadającą jej wielkość $C_{index,t}$ w łącznym indeksie strat (np. w indeksie PCS).

Dodatkowo w modelu określa się także odpowiednie procesy względem miernika wyceny o zneutralizowanym ryzyku.

W dalszej części artykułu zostanie przeprowadzona numeryczna analiza wyceny i zostaną omówione wyniki tej analizy¹⁷.

Model dynamiki aktywów

W typowym sposobie modelowania dynamiki aktywów zakłada się lognormalny rozkład składnika dyfuzji dla wartości aktywów, tak jak czynią to np. Merton i Cummins¹⁸. Główną wadą tego podejścia modelowego jest to, że nie pozwala ono uwzględnić explicite wpływu stochastycznych stóp procentowych na wartość aktywów. Jest to bardzo ważne dla modelowania wartości

¹⁶ Jin-Ping Lee, Min-Teh Yu, op. cit.

¹⁷ *Ryzyko zobowiązań i działań...*, op. cit.

¹⁸ J.D. Cummins, *Risk-based Premiums for Insurance Guarantee Funds*, „Journal of Finance” 1988, No. 43.

aktywów ubezpieczyciela, ponieważ ubezpieczyciele powszechnie posiadają w portfelu dużą część aktywów o stałym dochodzie. W szczególności ubezpieczyciele, którzy emitują obligacje CAT, inwestują głównie swe dochody ze sprzedanych obligacji CAT, w wysokiej jakości inwestycje wrażliwe na stopę procentową, takie jak weksle handlowe i papiery wartościowe skarbu państwa.

Okazuje się, że wyrażenie całkowitej wartości aktywów ubezpieczyciela w postaci dwóch składników ryzyka:

- ryzyka stopy procentowej,
- ryzyka kredytowego,

umożliwia pomiar wpływu ryzyka stopy procentowej na ceny obligacji CAT¹⁹.

Z teoretycznego punktu widzenia termin „ryzyko kredytowe”, o czym była mowa we wstępie artykułu, implikuje we wszystkich ryzykach ortogonalnych do ryzyka stopy procentowej.

W szczególności dynamikę wartości aktywów A_t ubezpieczyciela opisuje proces wyrażony równaniem stochastycznym (1), w którym wiodące znaczenie ma bieżący dryf μ_A będący pewną tendencją wynikającą z oddziaływania ryzyka kredytowego. Uwzględnia się w tym modelu także elastyczność ϕ bieżącej stopy procentowej aktywów ubezpieczyciela zmienności, wspomnianą już bieżącą stopę procentową r_t w czasie t , zmienność procesu ryzyka kredytowego σ_A oraz ryzyko kredytowe $W_{A,t}$ wyrażone procesem Wienera:

$$\frac{dA_t}{A_t} = \mu_A dt + \phi dr_t + \sigma_A dW_{A,t} \quad (1)$$

Model bieżącego oprocentowania

Model (1) przy założeniu, że bieżąca stopa procentowa modelowana jest zgodnie z procesem Coxa, Ingersolla i Rossa²⁰ wyrównywania pierwiastkowego (ang. *square-root*), tzn. *square-root* danej liczby, to liczba, jaka pomnożona przez siebie jest równa tej liczbie, co uniemożliwia pojawienie się ujemnej stopy procentowej, jaka może pojawić się w modelu Vasiceka²¹ i ma postać

$$dr_t = \kappa(m - r_t) dt + \nu \sqrt{r_t} dZ_t$$

gdzie κ jest uśrednionym wstecznym pomiarem wymuszonego oddziaływania; m jest przeciętną długoterminowej stopy procentowej; ν jest parametrem zmienności dla stopy procentowej; a Z_t jest procesem Wienera niezależnym od $W_{A,t}$, prowadzi do modelu (2) dynamiki aktywów:

¹⁹ J.C. Duan, A. Moreau, C.W. Sealey, *Deposit Insurance and Bank Interest Rate Risk: Pricing and Regulatory Implication*, „Journal of Banking and Finance” 1995, No. 19, s. 1091-1108.

²⁰ S.H. Cox, J. Ingersoll, S. Ross, *The Term Structure of Interest Rates*, „Econometrica” 1985, No. 53, s. 385-407.

²¹ O.A. Vasicek, *An Equilibrium Characterization of Term Structure*, „Journal of Financial Economics” 1977, No. 5, s. 177-188.

$$\frac{dA_t}{A_t} = (\mu_A + \phi \kappa m - \phi \kappa r_t) dt + \phi v \sqrt{r_t} dZ_t + \sigma_A dW_{A,t} \quad (2)$$

Dynamika aktywów ubezpieczyciela zneutralizowana względem ryzyka

Aby postać modelu (2) procesu dynamiki wartości aktywów była zneutralizowana względem ryzyka, należy zastosować mechanizm neutralizacji ryzyka. Dynamika dla procesu oprocentowania względem miernika wyceny o zneutralizowanym ryzyku, oznaczana przez Q , może być wyrażona równaniem

$$dr_t = \kappa^* (m^* - r_t) dt + v \sqrt{r_t} dZ_t^*$$

gdzie κ^* , m^* i Z^* są definiowane jako

$$\kappa^* = \kappa + \lambda_r$$

$$m^* = \frac{\kappa \cdot m}{\kappa + \lambda_r}$$

$$dZ_t^* = dZ_t + \frac{\lambda_r \cdot \sqrt{r_t}}{v} \cdot dt$$

Wielkość λ_r jest rynkową ceną ryzyka stopy procentowej i jest stałą względem procesu Coxa²²; Z_t^* jest procesem Wienera względem Q , której regułę można znaleźć także u Ritchkena²³.

Stąd, dynamika aktywów ubezpieczyciela zneutralizowana względem ryzyka wyraża równanie

$$\frac{dA_t}{A_t} = r_t dt + \phi v \sqrt{r_t} dZ_t^* + \sigma_A dW_t^* \quad (3)$$

gdzie W^* jest procesem Wienera względem Q , niezależnym od procesu Z_t^* .

Dynamika łącznej straty

Stosując typowe w analizie aktuarialnej podejście ustalania dynamiki strat²⁴, model łącznej straty wyrażony może zostać w postaci złożonego procesu Poissona, sumy nagłych wzrostów (ang. *a sum of jumps*).

²² S.H. Cox, J. Ingersoll, S. Ross, op. cit.

²³ P. Ritchken, *Derivative Markets: Theory, Strategy, and Applications*, Harper Collins, New York 1996.

²⁴ N.L. Bowers, H.U. Gerber, J.C. Hickman et. al., *Actuarial Mathematics*, The Society of Actuaries, Itasca, Illinois 1986.

W celu oszacowania wpływu ryzyka stopy bazowej na cenę obligacji CAT wprowadza się wielkość $C_{i,t}$ oznaczającą łączną stratę dla firmy i emitującej obligację, a także odpowiadającą jej wielkość $C_{index,t}$ w łącznym indeksie strat (np. w indeksie PCS). Te dwa procesy mogą być opisane następująco:

$$C_{i,t} = \sum_{j=1}^{N(t)} X_{i,j} \quad (4)$$

i

$$C_{index,t} = \sum_{j=1}^{N(t)} X_{index,j} \quad (5)$$

gdzie proces $\{N(t)\}_{t \geq 0}$ jest procesem liczby strat opisanym przez proces Poissona z intensywnością λ . Symbole $X_{i,j}$ i $X_{index,j}$ oznaczają odpowiednio łączną stratę spowodowaną przez katastrofę j w trakcie danego okresu dla firmy ubezpieczeniowej emitującej obligację CAT i łączny indeks strat. Zakłada się, że wielkości $X_{i,j}$ ($X_{index,j}$) dla $j = 1, 2, \dots, N(T)$, są wzajemnie niezależne i mają taki sam rozkład lognormalny oraz są także niezależne od procesu liczby straty, a ich logarytmiczne średnie i wariancje są odpowiednio oznaczone przez μ_i (μ_{index}) i σ_i^2 (σ_{index}^2). Ponadto zakłada się, że współczynniki korelacji ρ logarytmów względem $X_{i,j}$ ($X_{index,j}$) dla różnych $j = 1, 2, \dots, N(T)$, są jednakowe.

Dynamika straty względem miernika wyceny o zneutralizowanym ryzyku

Aby dokonać wyceny obligacji CAT, należy znać dynamikę straty względem miernika wyceny o zneutralizowanym ryzyku Q . Kiedy proces straty ma nagle wzrosty, rynek nazywany jest wtedy niezupełnym i nie ma unikalnego miernika wyceny. A zatem, postępując analogicznie jak Merton²⁵, zakłada się, że warunki ekonomiczne wyznaczone są tylko w sensie marginalnym (krajcowym) pod wpływem zlokalizowanych katastrof, takich jak trzęsienia ziemi i huragany, oraz że proces liczby strat $\{N(t)\}$ i wielkość strat $X_{i,j}$ ($X_{index,j}$), są bezpośrednio związane z idiosynkratycznymi „szokami” dla rynków kapitałowych, czyli czynnikami wpływającymi na rynki kapitałowe w sposób niepożądany.

Czynniki takie – szoki katastrofalne reprezentują ryzyko „niesystematyczne” i odpowiadają im zerowe stopy składki za ryzyko (ang. *risk premium*), którego są generatorami.

²⁵ R. Merton, *An Analytic Derivation of the Cost of Deposit Insurance and Loan Guarantee*, „Journal of Banking and Finance” 1977, No. 1, s. 3-11.

Zakładając, że takie ryzyko wzrostu jest niesystematyczne i możliwe do zróżnicowania (dywersyfikacji), dodanie składki za to ryzyko jest niepotrzebne. Okazuje się, że założenie to jest ważne, ponieważ nie można zastosować oceny o neutralnym ryzyku do sytuacji, w których wielkość wzrostu jest systematyczna. Kwestię tę szczegółowo omawiają Naik i Lee, Cummins i Geman, Cox i Pedersen²⁶.

Można zatem przyjąć, że procesy łącznych strat wyrażone w równaniach (4) i (5), zachowują swe oryginalne charakterystyki rozkładu po zmianie z fizycznego miernika prawdopodobieństwa na miernik wyceny o zneutralizowanym ryzyku.

4. Wpływ ryzyka stopy bazowej i moralnego hazardu na cenę oraz wypłatę obligacji CAT

Okazuje się, że oszacowanie ceny obligacji CAT jest możliwe, gdy znane są procesy o neutralnym ryzyku dynamiki aktywów, straty i stopy procentowej, poprzez dyskontowanie przewidywania ich różnych wypłat w środowisku o neutralnym ryzyku. Natomiast określenie wypłaty obligacji CAT można przeprowadzić przy alternatywnych założeniach dotyczących ryzyka wypłaty. Rozważyć można przypadek podstawowy, gdy nie występuje ryzyko odmowy zapłaty²⁷, a także przypadek obligacji uwzględniających w swych warunkach ryzyko odmowy zapłaty z potencjalnym ryzykiem stopy bazowej i moralny hazard.

Obligacje CAT w warunkach niewystępowania ryzyka odmowy wypłaty

Dla wyceny obligacji CAT przyjmuje się, że jest to hipotetyczna dyskontowana obligacja, z której płatności (PO_T) w terminie płatności (tzn. czasie T) są następujące:

$$PO_T = \begin{cases} a \cdot L & \text{gdy } C_T \leq K \\ r \cdot p \cdot a \cdot L & \text{gdy } C_T > K \end{cases} \quad (6)$$

²⁶ Por. V. Naik, M. Lee, *General Equilibrium Pricing of Options on the Market Portfolio With Discontinuous Returns*, „Review of Financial Studies” 1990, No. 3(4), s. 493-521; J.D. Cummins, H. Geman, op. cit.; S.H. Cox, H.W. Pedersen, *Catastrophe Risk Bonds*, „North American Actuarial Journal” 2000, No. 4(4), s. 56-82.

²⁷ Ryzyko zobowiązań i działań..., op. cit.

gdzie:

- K – poziom uruchomienia wykupu obligacji ustalony w postanowieniach obligacji CAT bond;
- C_T – łączna strata w terminie płatności;
- $r \cdot p$ – część kapitału, jaka musi być wypłacona posiadaczom obligacji, kiedy nastąpiło uruchomienie umorzenia;
- L – nominalna wartość całkowitych długów firmy emitującej, której część stanowi nominalna wartość obligacji CAT;
- a – stosunek nominalnej sumy obligacji CAT do całkowitych należnych długów.

Wycena obligacji CAT z płatnościami określonymi w równaniu (6) przeprowadzana jest przy założeniu, że zmienne strukturalne θ i η , które determinują strukturę terminową stopy procentowej i straty łącznej, wystarczająco dobrze „zachowują się”, aby móc zastosować podejście o neutralnym ryzyku Coxa i Rossa²⁸ i Harrisona i Pliska²⁹. Dokładniej określając, przy mierniku wyceny o zneutralizowanym ryzyku Q , cena obligacji CAT w dniu emisji (tzn. momencie 0) może być wyrażona przez wielkość $E_{\theta, \eta}^*$ oznaczającą jej oczekiwaną wartość w środowisku o neutralnym ryzyku.

W dalszej specyfikacji modelu wyceny zakłada się, że zmienne strukturalne θ , które dla celu oszacowania obligacji ryzyka katastroficznego określają strukturę czasową stóp procentowych, są niezależne od zmiennych strukturalnych η , które odnoszą się do zmiennych ryzyka katastrofalnego. Przy tych założeniach cenę obligacji CAT zależną od czynnika ceny oznaczonej przez $B_{CIR}(\theta, T)$, dotyczącej obligacji wolnej od ryzyka, który można znaleźć w literaturze³⁰, można wyrazić następująco:

$$P_{CAT}(0) = B_{CIR}(0, T) \cdot \sum_{j=0}^{\infty} \exp(-\lambda \cdot T) \cdot \frac{(\lambda \cdot T)^j}{j!} \cdot F^j(K) + r \rho \cdot \left(1 - \sum_{j=0}^{\infty} \exp(-\lambda \cdot T) \cdot \frac{(\lambda \cdot T)^j}{j!} \cdot F^j(K)\right) \quad (7)$$

gdzie:

$$F^j(K) = P(X_{i,1} + X_{i,2} + \dots + X_{i,j} \leq K), \quad F^j \text{ splot } j \text{ dystrybuant,}$$

$$B_{CIR}(0, T) = A(0, T) \cdot \exp[-B(0, T) \cdot r],$$

²⁸ S.H. Cox, S. Ross, op. cit.

²⁹ J.M. Harrison, S.R. Pliska, *Martingales and Stochastic Integrals in the Theory of Continuous Trading*, „Stochastic Processes and Their Applications” 1981, No. 11, s. 215-260.

³⁰ S.H. Cox, J. Ingersoll, S. Ross, op. cit.

gdzie:

$$A(0, T) = \left[\frac{2 \cdot \gamma \cdot \exp\left(\kappa + \gamma \cdot \frac{T}{2}\right)}{(\kappa + \gamma) \cdot \exp(\gamma^T - 1) + 2 \cdot \gamma} \right]^{\frac{2 \cdot \kappa \cdot m}{v^2}}$$

$$B(0, T) = \frac{2 \cdot [\exp(\gamma^T) - 1]}{(\kappa + \gamma) \cdot [\exp(\gamma^T) - 1] + 2 \cdot \gamma}$$

$$\gamma = \sqrt{\kappa^2 + 2 \cdot v^2}$$

Aproksymacja rozkładu łącznej straty i ceny obligacji – analityczne rozwiązanie

Przy założeniu, że składniki wielkości straty katastrofalnej są niezależne, o tym samym rozkładzie lognormalnym, dokładny rozkład łącznej straty w terminie płatności, oznaczony przez $f(C_T)$, nie jest oczywiście znany w dokładnej postaci. Można jednak wyznaczyć przybliżoną analityczną postać tego rozkładu prawdopodobieństwa. W tym celu aproksymuje się dokładną dystrybuantę przez dystrybuantę lognormalną, oznaczoną przez $g(C_T)$, o określonych momentach. Podobnie postąpili Jarrow i Rudd, Turnbull i Wakeman oraz Nielson i Sandmann, wykorzystując te same założenia w przybliżaniu wartości opcji azjatyckich i tzw. opcji koszykowych³¹.

Warunki w zastosowaniu tego podejścia wymagają jedynie ustalenia pierwszych dwóch momentów rozkładu wyznaczonego przez funkcję $g(C_T)$, jako równych momentom dokładnego (lecz nieznanego) rozkładu łącznej straty płatności $f(C_T)$. Zanotujemy to w postaci:

$$\text{moment rzędu pierwszego} \quad \mu_g = E[C] = \lambda \cdot T \cdot \exp\left\{\mu_i + \frac{1}{2} \cdot \sigma_i^2\right\} \quad (8)$$

$$\text{moment centralny rzędu drugiego} \quad \sigma_g = \text{Var}[C] = \lambda \cdot T \cdot \exp\{2 \cdot \mu_i + 2 \cdot \sigma_i^2\} \quad (9)$$

a zatem μ_g i σ_g^2 oznaczają odpowiednio średnią i wariancję aproksymującego rozkładu $g(C)$.

³¹ Por. R. Jarrow, A. Rudd, *Approximate Option Valuation for Arbitrary Stochastic Processes*, „Journal of Financial Economics” 1982, No. 10, s. 347-369; S. Turnbull, L. Wakeman, *A Quick Algorithm for Pricing European Average Option*, „Journal of Financial and Quantitative Analysis” 1991, No. 26, s. 377-389; J. Nielson, K. Sandmann, *The Pricing of Asian Options Under Stochastic Interest Rate*, „Applied Mathematical Finance” 1996, No. 3, s. 209-236.

Wartość obligacji CAT może więc zostać wyrażona wzorem:

$$B_{apr}(0) = B_{CIR}(0, T) \cdot \int_0^K \frac{1}{\sqrt{2 \cdot \pi \cdot \sigma_g \cdot C_T}} \cdot \exp\left\{-\frac{1}{2} \cdot (\ln C_T - \mu_g)^2\right\} dC_T + \\ + rp \cdot \int_K^\infty \frac{1}{\sqrt{2 \cdot \pi \cdot \sigma_g \cdot C_T}} \cdot \exp\left\{-\frac{1}{2} \cdot (\ln C_T - \mu_g)^2\right\} dC_T \quad (10)$$

gdzie $B_{apr}(0)$ jest przybliżoną analityczną wyceną obligacji CAT w momencie 0. Reguła ta jest podobna do wyprowadzonej przez Litzenbergera, Beaglehole'a i Reynoldsa, z tym że w ich modelu stopa procentowa jest stała.

W końcowej części empirycznej tego artykułu analityczne rozwiązanie zostanie porównane z oszacowaniami wyprowadzonymi metodą numeryczną, w których nie zakłada się przyjmowanych tu przybliżeń.

Obligacje CAT w warunkach ryzyka niewypłacalności ubezpieczyciela

Założenie braku odmowy wypłaty przyjmowane wyżej odnosi się do genezy analitycznej prezentacji obligacji CAT. W przypadku praktycznych rozważań dotyczących ryzyka odmowy wypłaty przez ubezpieczyciela, ryzyka stopy bazowej i moralnego hazardu odnoszących się do obligacji CAT, zostaną określone ich wypłaty, a następnie zostanie dokonana ich ocena z zastosowaniem metody numerycznej, co zostanie zaprezentowane w części końcowej rozdziału.

W przypadku gdy ubezpieczyciel staje się niewypłacalny i odmawia wypłaty, założymy, że posiadacze obligacji CAT mają priorytet przejawiający się tym, że zabezpieczają w pierwszej kolejności innych posiadaczy długów, ponieważ dochody ze sprzedaży obligacji CAT są zwykle inwestowane w fundusz powierniczy i mogą być zlikwidowane jedynie w celu wypłaty ograniczonych roszczeń lub zwrotów dla posiadaczy obligacji. Jeśli ubezpieczyciel jest wypłacalny, wówczas posiadacze obligacji CAT mogą otrzymać pełen kapitał obligacji CAT, pod warunkiem, że bazowe straty są niższe niż poziom uruchomienia wypłat; w innym przypadku może być zwrócona jedynie część kapitału.

Określenie wypłat przy braku ryzyka stopy bazowej

W pierwszym rozważanym przypadku będą określone wypłaty w warunkach ryzyka niewypłacalności ubezpieczyciela, gdy nie ma ryzyka stopy bazowej. Oznacza to, że dług ubezpieczyciela jest umarzany, kiedy jego rzeczywista strata jest większa niż ustalona wielkość straty w warunkach umowy.

Ryzyko bazowe odnosi się do ryzyka pojawiającego się wtedy, gdy straty jakie ponoszą poszczególni ubezpieczyciele nie będą antycypowane z założonym ich skorelowaniem z bazowym indeksem straty obligacji CAT. Ryzyko stopy bazowej może redukować wpływ zabezpieczenia obligacji CAT i zwiększyć prawdopodobieństwo odmowy wypłaty przez firmę emitującą obligacje.

Wypłaty w przypadku braku ryzyka bazowego w warunkach ryzyka odmowy wypłaty obligacji CAT mogą być zapisane następująco

$$PO_{i,T} = \begin{cases} a \cdot L, & \text{gdy } C_{i,T} \leq K \text{ oraz } C_{i,T} \leq A_{i,T} - a \cdot L \\ rp \cdot a \cdot L, & \text{gdy } K < C_{i,T} \leq A_{i,T} - rp \cdot a \cdot L \\ \max\{A_{i,T} - C_{i,T}, 0\} & \text{w innych przypadkach} \end{cases} \quad (11)$$

gdzie:

$PO_{i,T}$ – wypłaty w terminie płatności dla obligacji CAT umorzone względem rzeczywistych strat firmy emitującej;

$A_{i,T}$ – wartość aktywów w terminie płatności firmy emitującej;

$C_{i,T}$ – łączna strata firmy emitującej obligacje w terminie płatności;

L , a , K oraz rp są zdefiniowane tak jak we wcześniej omówionych przypadkach.

Określenie wypłat w warunkach ryzyka stopy bazowej

W tym przypadku porównywane są wypłaty z wypłatami, gdy występuje ryzyko stopy bazowej, w których dług jest umarzany przy pewnym złożonym indeksie strat.

W przypadku obligacji CAT umarzanej względem złożonego indeksu strat, wypłaty w warunkach ryzyka niewypłacalności ubezpieczyciela mogą być wyrażone następująco:

$$PO_{indeks,T} = \begin{cases} a \cdot L, & \text{gdy } C_{indeks,T} \leq K \text{ oraz } C_{i,T} \leq A_{i,T} - a \cdot L \\ rp \cdot a \cdot L, & \text{gdy } C_{indeks,T} > K \text{ oraz } K < C_{i,T} \leq A_{i,T} - rp \cdot a \cdot L \\ \max\{A_{i,T} - C_{i,T}, 0\} & \text{w innych przypadkach} \end{cases} \quad (12)$$

gdzie:

$C_{indeks,T}$ – wartość złożonego indeksu w terminie płatności;

a , L , rp , $A_{i,T}$, $C_{i,T}$ i K oznaczają te same wielkości, co w równaniu (13).

Ze struktury wypłat z takich obligacji CAT i określonej dynamiki aktywów oraz stopy procentowej, obligacje CAT mogą być oszacowane następująco:

$$P_i(0) = \frac{1}{a \cdot L} \cdot E_0^*[\exp\{-\bar{r} \cdot T \cdot PO_{i,T}\}] \quad (13)$$

gdzie:

$P_i(0)$ – cena obligacji CAT w warunkach ryzyka odmowy wypłaty przy braku ryzyka stopy bazowej;

E_0^* – oczekiwania założone w dniu emisji przy mierniku wyceny Q ;

$\bar{r}$ – przeciętna stopa procentowa wolna od ryzyka pomiędzy momentem emisji a terminem płatności;

$\frac{1}{a \cdot L}$ – współczynnik normalizujący ceny obligacji CAT dla nominalnej wartości 1 dolara,

oraz

$$P_{index}(0) = \frac{1}{a \cdot L} \cdot E_0^*[\exp\{-\bar{r} \cdot T \cdot PO_{index,T}\}] \quad (14)$$

gdzie:

$P_{index}(0)$ – wartość ceny obligacji CAT w warunkach ryzyka odmowy wypłaty z ryzykiem stopy bazowej w momencie 0;

$E_0^*, \bar{r}, \frac{1}{a \cdot L}$ – zdefiniowane są tak jak we wzorze (12).

Określenie wypłat z uwzględnieniem moralnego hazardu

Trzeci z możliwych wariantów wyceny obligacji w warunkach ryzyka niewypłacalności ubezpieczyciela dotyczy struktury modelowania wypłat w warunkach istnienia moralnego hazardu. Ten przypadek zostanie w dalszej części dokładniej przeanalizowany.

Z taką sytuacją możemy mieć do czynienia wtedy, gdy obligacja CAT jest umorzona względem strat firmy emitującej, ponieważ firma emitująca ma priorytet i inicjuje rozliczenia roszczeń oraz, bardziej „hojnie” rozlicza roszczenia, kiedy poniesiona strata zbliża się do poziomu uruchomienia wypłat z obligacji. W tym modelu zakłada się, że firma emitująca obligacje „łagodzi” warunki skonstruowanej polisy rozliczeniowej, gdy skumulowane straty spadają do zakresu zamknięcia uruchomienia (ang. *trigger*) wypłat. Założenie to jest uzasadnione tym, że kumulacja strat spowodowałaby wówczas wzrost oczekiwanych strat dla kolejnego wydarzenia o charakterze katastrofalnym.

W tym przypadku zmianę procesu straty należy określić następująco:

$$\mu'_i = \begin{cases} (1 + \alpha) \cdot \mu_i, & \text{gdy } (1 - \beta) \cdot K \leq C_{i,j} \leq K \\ \mu_i & \text{w innym przypadku} \end{cases} \quad (15)$$

gdzie μ' jest logarytmiczną średnią strat poniesionych w wyniku wystąpienia katastrofy zidentyfikowanej numerem $(j + 1)$, gdy skumulowana wartość straty $C_{i,j}$ spada do określonego zakresu, $(1 - \beta) K \leq C_{i,j} \leq K$; α jest dodatnią stałą, która jest interpretowana (lub odpowiada) jako procentowy wzrost średniej; β jest dodatnią stałą, specyfikującą zakres występowania hazardu moralnego lub sposób postępowania, który nie musi być bliżej określony, a który wprowadza się tu *ad hoc*, w celu ustalenia podejścia do hazardu moralnego, gdy jest taka potrzeba, ale bez wcześniejszego planowania występowania zespołu uwarunkowań prowadzących do jego zaistnienia. W kontekście tego ostatniego założenia dotyczącego stałej β , a modyfikującego zmianę procesu straty, należy zaznaczyć, że w literaturze nie rozstrzyga się tego aspektu hazardu moralnego i należy go zatem postrzegać raczej jako wyzwanie inspirujące innych badaczy do modelowania hazardu moralnego.

Przedstawiona struktura analityczna kompleksowego kontraktu kontyngentowego (ang. *complex contingent contract*) jest podstawą do jej uzupełnień i modyfikacji, ale przedstawiona w obecnej postaci umożliwia ona jednak numeryczne szacowanie ceny obligacji CAT. Jak już sygnalizowaliśmy, będzie to przedmiotem egzemplifikacji zastosowania modelu szacowania obligacji CAT na danych symulowanych z zastosowaniem metody Monte Carlo.

5. Egzemplifikacja struktury modelowej obligacji CAT z zastosowaniem symulacji Monte Carlo

Dla wyjaśnienia dość trudnej z formalnego punktu widzenia struktury analitycznej szacowania obligacji CAT zostanie przeprowadzona egzemplifikacja aplikacji struktury modelowej szacowania tej obligacji³². Ceny obligacji CAT były szacowane³³ z zastosowaniem symulacji Monte Carlo. Zostanie rozpatrzona wycena obligacji CAT w różnych konfiguracjach uwzględniających lub nie hazard moralny, ryzyko bazowe i ryzyko odmowy zapłaty.

³² Jin-Ping Lee, Min-teh Yu, op. cit., *Ryzyko zobowiązań i działań...*, op. cit.

³³ Jin-Ping Lee, Min-teh Yu, op. cit.

Punktem wyjścia w szacowaniu ceny obligacji CAT będzie ustalony zbiór parametrów i wartości bazowych. Natomiast w celu oszacowania porównawczego wpływu tych parametrów na ceny obligacji CAT też są ustalone odchylenia od wartości bazowych. Dla uproszczenia zakłada się, że **całkowita kwota długów firmy emitującej, która jest wyrażona w obligacjach CAT, ma wartość nominalną 100 \$, a termin wygaśnięcia-płatności obligacji CAT jest równy jeden rok.** Symulacje zostały uruchomione na tygodniowej bazie danych z 20 000 ścieżkami (*ang. paths*). Dane parametry i wartości bazowe zawarte są w tabeli 1.

Tabela 1

Parametry i wartości bazowe

Rodzaj parametrów i wartości bazowych	Wartości
Parametry aktywów	
A – aktywa ubezpieczyciela	Aktywa do zobowiązań A/L : 1,1, 1,2 i 1,3
μ_A – tendencja spowodowana ryzykiem kredytowym (<i>drift due</i>)	
Φ – elastyczność stopy procentowej aktywów	0, -3, -5
σ_A – zmienność ryzyka kredytowego	5%
W_A – winerowski szok kredytowy	

Źródło: Na podstawie: Jin-Ping Lee, Min-teh Yu, *Pricing Default-Risky CAT Bonds with Moral Hazard and Basis Risk*, „Journal of Risk and Insurance” 2002, Vol. 69.

Współczynniki początkowej pozycji aktywów/ zobowiązań (lub kapitału) (A/L) są ustalone odpowiednio jako 1.1, 1.2 i 1.3. Przeciętny współczynnik A/L dla sektora ubezpieczeń wynosił w USA około 1.3, na podstawie wartości księgowej, przez ostatnie dziesięć lat. Elastyczność stopy procentowej aktywów ubezpieczyciela jest ustalona na 0, -3 i -5, i odpowiednio do pomiaru odzwierciedla wpływ ryzyka stopy procentowej ubezpieczyciela ceny obligacji CAT. Zmienność zwrotów z aktywów, jaka jest spowodowana przez ryzyko kredytowe, jest ustalona na poziomie 5%.

W tabeli 2 są podane parametry stopy procentowej. Początkowa (*ang. the initial spot*) stopa procentowa r i długookresowa stopa procentowa m są ustalone na poziomie 5%. Wartość średniego wstecznego oddziaływania κ (*ang. magnitude of mean-reverting force*), tzn. o charakterze sprzężenia zwrotnego, jest ustalona jako 0,2, podczas gdy zmienność v stopy procentowej jest ustalona na 10%. Ceny λ_r rynkowe ryzyka stopy procentowej są odpo-

wiednio ustalone jako wartości 0 i $-0,01$. Wszystkie te parametry struktury czasowej są zawarte w zakresie typowo wykorzystywanym (*previous*) w literaturze.

Tabela 2

Parametry stopy procentowej

Rodzaj parametrów	Wartości
Parametry stopy procentowej	
r – początkowa natychmiastowa stopa procentowa	5%
κ – wielkość średniej-powracającej siły	0,2
M – długookresowa średnia stopa procentowa	5%
zmienność v stopy procentowej	10%
ceny λ_r rynkowe ryzyka stopy procentowej	0, $-0,01$
Z – winerowski proces szokowej stopy procentowej	

Źródło: Na podstawie: Jin-Ping Lee, Min-teh Yu, op. cit.

W tabeli 3 są zamieszczone parametry straty katastrofalnej i inne parametry, w tym poziomy uruchomienia umorzenia długu i współczynnik kapitału odpowiadający wypłacie, jeśli umorzenie długu będzie uruchomione.

Intensywność pojawienia się strat katastrofalnych jest ustalona odpowiednio na trzech poziomach 0,5, 1 i 2, które odzwierciedlają częstotliwości zdarzeń katastrofalnych zdarzających się w jednym roku. Zakładamy także, że wartości parametru dla straty katastrofalnej, jako takie same dla poszczególnych (indywidualnych) ubezpieczycieli i łącznego (*the composite*) wskaźnika straty. Średnie logarytmiczne μ_i i μ_{indeks} ustala się na poziomie 2, a logarytmiczne odchylenia standardowe σ_i i σ_{indeks} na 0,5, 1 i 2. Wartości dla wskaźnika i poszczególnych ubezpieczycieli może być oczywiście rozmaicie modyfikowane, ale powiększa to złożoność numeryczną obliczeń, nie powodując jednocześnie poszerzenia analizy ryzyka stopy bazowej. Analiza skupi się raczej na współczynniku korelacji ρ_x pomiędzy poszczególną stratą a wskaźnikiem straty, niż na ich średnich i odchyleniach standardowych. Część kapitału rp konieczna przy wypłatach roszczeń w momencie uruchomienia umorzenia długu jest ustalona i wynosi 0,5. Stosunek wartości obligacji CAT do całkowitego długu ubezpieczyciela a jest ustalony na 0,1. Dodatkowo ustalone są trzy różne poziomy uruchomienia K umorzenia długu: 100, 110 i 120.

Tabela 3

Parametry straty katastrofalnej

Rodzaj parametrów	Wartości
Parametry straty katastrofalnej	
μ_i – średnia logarytmu ilości strat katastroficznych dla ubezpieczyciela	2
μ_{indeks} – średnia logarytmu ilości strat katastroficznych dla złożonego wskaźnika straty	0,2
σ_i – odchylenie standardowe logarytmu wartości strat katastroficznych ubezpieczyciela	0,5, 1, 2
σ_{indeks} – odchylenie standardowe logarytmu wartości strat katastroficznych ubezpieczyciela dla łącznego wskaźnika straty	0,5, 1, 2
ρ_x – współczynnik korelacji logarytmów wartości strat katastroficznych ubezpieczyciela i łącznego straty 0,5, 0,8, 1	0,5, 0,8, 1
$N(t)$ – proces Poissona modelujący dynamikę pojawienia się katastrof	
Inne parametry	
K – poziomy uruchomienia	100, 110, 120
R_p – współczynnik kapitału jaki musi być zapłacony, jeśli umorzenie długu jest uruchomione	0,5
a – stosunek kwoty obligacji CAT do całkowitych długów	0.1
α – intensywność moralnego hazardu	20%
β – współczynnik ustalony poniżej uruchomienia, jaki spowoduje moralny hazard ubezpieczyciela	20%
L – całkowita ilość długów ubezpieczyciela	100

Źródło: Na podstawie: Jin-Ping Lee, Min-teh Yu, op. cit.

5.1. Numeryczne wyznaczenie cen obligacji CAT w warunkach ryzyka odmowy wypłaty z alternatywnym uwzględnieniem występowania hazardu moralnego

Tabela 4 zawiera ceny obligacji w warunkach ryzyka odmowy wypłaty z alternatywnym uwzględnieniem niewystępowania moralnego hazardu i występowaniem hazardu moralnego przy alternatywnych wartościach początkowej pozycji kapitału (A/L), intensywności katastrofy, zmiennej straty i elastyczności stopy procentowej aktywów firmy emitującej. Trzy (górna, środkowa i dolna) wartości relacjonowane w każdej z komórek tabeli 4 reprezentują odpowiednie

oszacowanie przy elastyczności stopy procentowej równej odpowiednio 0, -3 i -5. Wyższa (absolutna) wartość elastyczności stopy procentowej odpowiada wyższej zmienności aktywów i ryzyku odmowy wypłaty przez emitenta. A zatem oczekiwane jest, że górna wartość każdej komórki (cena obligacji CAT dla $\phi = 0$) będzie wyższa niż środkowa (cena obligacji CAT dla $\phi = -3$) i dolna wartość (cena obligacji CAT dla $\phi = -5$).

Oczekiwane jest także, że im wyższa jest początkowa pozycja kapitału (A/L) firmy emitującej, tym niższe ryzyko odmowy wypłaty i wyższe ceny obligacji CAT. Wzrost ceny obligacji jest spowodowany intensywnością zdarzenia i zmiennością straty katastrofalnej. Ceny obligacji są określone względem nominalnej kwoty jednego dolara. Górna wartość, środkowa wartość i dolna wartość w każdej komórce to ceny obligacji CAT, obliczone gdy elastyczności stopy procentowej aktywów wynoszą odpowiednio 0, -3, -5.

Tabela 4

Ceny obligacji CAT w warunkach ryzyka niewypłacalności przy braku i z moralnym hazardem ($\lambda_r = -0,01$)

(λ, σ)	$A/L = 1,1$			$A/L = 1,2$			$A/L = 1,3$		
	<i>Uruchomienia</i>								
	100	110	120	100	110	120	100	110	120
	<i>Brak moralnego hazardu</i>								
(0,5, 1)	0,94891	0,94919	0,94924	0,94926	0,94968	0,95002	0,94940	0,94993	0,95021
	0,94864	0,94898	0,94898	0,94910	0,94958	0,94963	0,94931	0,94984	0,95002
	0,94847	0,94876	0,97876	0,94896	0,94944	0,94953	0,94977	0,95029	0,95062
(2,1)	0,93121	0,93292	0,93335	0,93360	0,93688	0,93845	0,93479	0,93827	0,94046
	0,93078	0,93282	0,93335	0,93321	0,93611	0,93769	0,93474	0,93821	0,94021
	0,92959	0,93150	0,93231	0,93268	0,93563	0,93706	0,93442	0,93761	0,93942
(2,2)	0,75871	0,76281	0,76323	0,76500	0,77104	0,77470	0,77044	0,77696	0,78314
	0,75743	0,76095	0,76191	0,76476	0,77075	0,77446	0,76974	0,77621	0,78211
	0,75658	0,76024	0,76229	0,76317	0,76317	0,77212	0,76955	0,77554	0,78063
	<i>Występowanie moralnego hazardu</i>								
	$A/L = 1,1$			$A/L = 1,2$			$A/L = 1,3$		
	100	110	120	100	110	120	100	110	120
(0,5, 1)	0,94695	0,94738	0,94743	0,94746	0,94817	0,94841	0,94781	0,94853	0,94881
	0,94667	0,94705	0,94709	0,94726	0,94778	0,94792	0,94767	0,94838	0,94867
	0,94634	0,94672	0,94677	0,94712	0,94774	0,94788	0,94757	0,94824	0,94847
(1,2)	0,84697	0,84892	0,84925	0,85073	0,85444	0,85615	0,85394	0,85770	0,85079
	0,84587	0,84759	0,84831	0,85054	0,85363	0,85525	0,85384	0,85750	0,86050
	0,84551	0,84727	0,84727	0,84988	0,85268	0,85426	0,85369	0,85698	0,85936
(2,2)	0,72243	0,72628	0,72690	0,72928	0,73584	0,74017	0,73542	0,74226	0,74883
	0,72064	0,72416	0,72521	0,72877	0,73463	0,73853	0,73456	0,74117	0,74702
	0,71921	0,72254	0,72440	0,72764	0,73278	0,73645	0,73378	0,74001	0,74539

Źródło: Dane symulowane opracowane na podstawie Jin-Ping Lee, Min-teh Yu, op. cit.

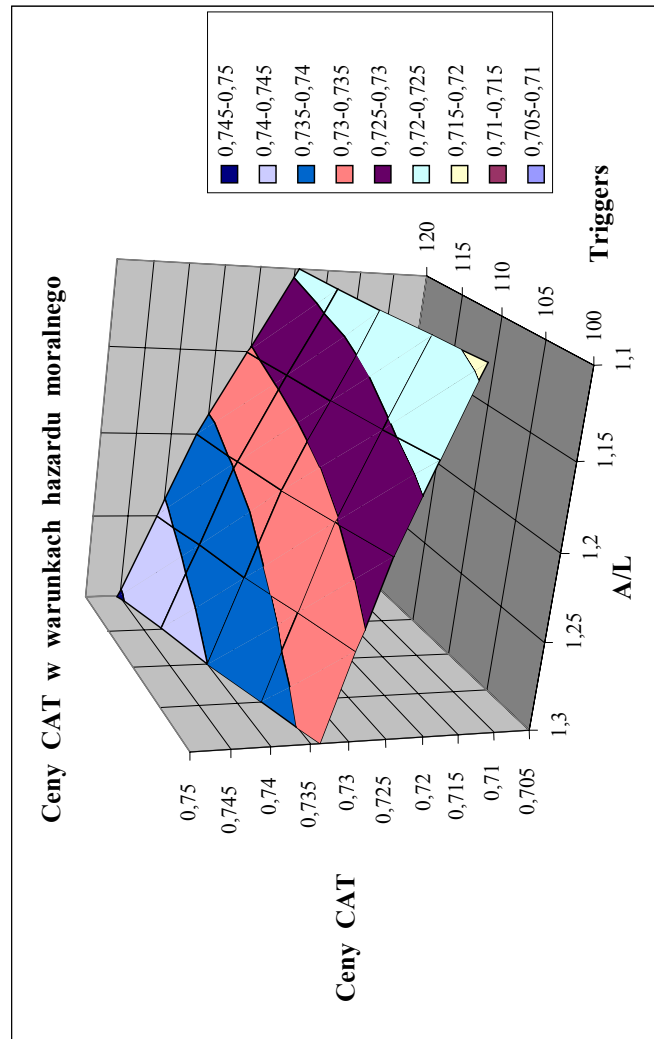
Zaobserwujemy, że składka za ryzyko odmowy wypłaty zmniejsza się wraz ze współczynnikiem A/L i wzrasta wraz z intensywnością zdarzenia i zmiennością straty. Składka na ryzyko odmowy zapłaty może wzrosnąć do 1,015 punktów bazowych dla przypadku $A/L = 1,1$, $(\lambda, \sigma) = (2, 2)$ i $\phi = -5$.

Aby uwzględnić wpływ moralnego hazardu, zakłada się, że gdy nagromadzone straty sięgają 80% poziomu uruchomienia, ubezpieczyciel rozliczy roszczenia katastrofalne bardziej hojnie, co spowoduje zwiększenie oczekiwanej straty katastroficznej o 20%, co odzwierciedla współczynnik intensywności moralnego hazardu $\alpha = 0,2$. Moralny hazard zatem zwiększa ryzyko odmowy wypłaty i obniża wartość obligacji. Na przykład, gdy $(\lambda, \sigma) = (2, 2)$, $\phi = -5$, cena spada o około 350 punktów bazowych. Wielkość wpływu moralnego hazardu wzrasta wraz z (λ, σ_i) i absolutną wartością ϕ , oraz zmniejsza się ze współczynnikiem A/L .

Kształtowanie się ceny obligacji CAT w warunkach ryzyka niewypłacalności bez moralnego hazardu $(\lambda, \sigma) = (2, 2)$ i $\phi = 0$, przy różnych poziomach uruchomienia (*triggers*) A/L wypłaty obligacji przedstawione zostało w pracy W. Szkutnika³⁴. Stwierdzono wzrastanie ceny obligacji CAT wraz ze wzrostem uruchomienia. Na wzrost ten miał także wpływ wzrost intensywności λ i poziom zmienności straty σ_i .

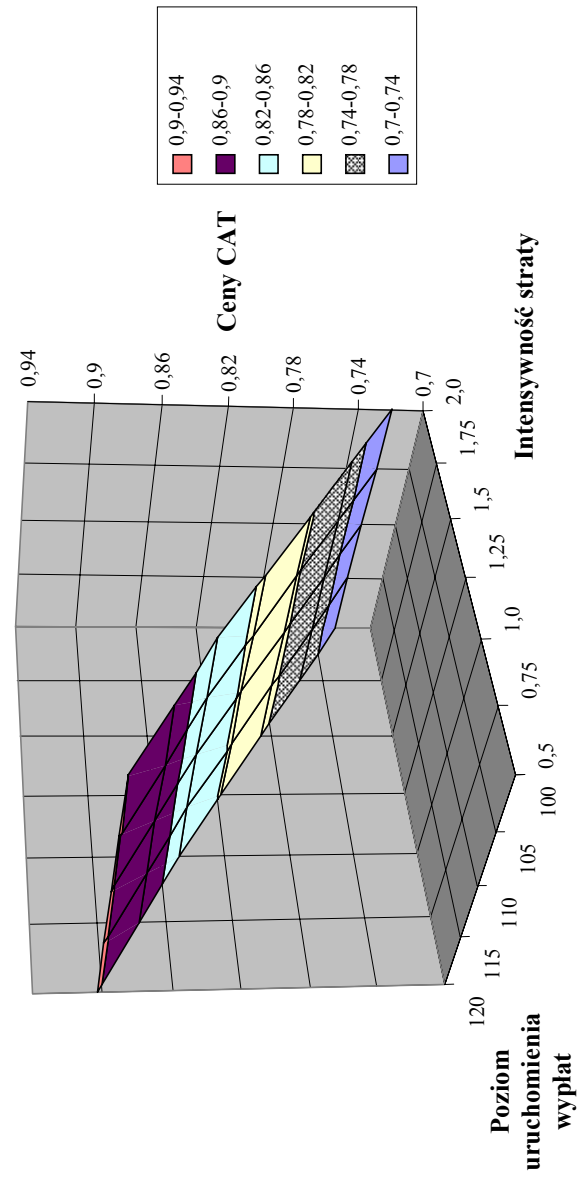
Na rysunku 1 przedstawione jest kształtowanie się ceny obligacji CAT w warunkach ryzyka niewypłacalności z uwzględnieniem występowania moralnego hazardu oraz parametrami intensywności, zmienności i elastyczności: $(\lambda, \sigma) = (2, 2)$, i $\phi = -5$, przy różnych poziomach uruchomienia (*triggers*) A/L wypłaty obligacji. Na osi odciętych zaznaczone są poziomy uruchomienia: 1,1, 1,15, 1,2, 1,25, 1,3. Widoczne jest także wzrastanie ceny obligacji CAT wraz ze wzrostem uruchomienia. Na wzrost ten ma oczywiście także wpływ wzrost intensywności λ i poziom zmienności straty σ_i .

³⁴ W. Szkutnik, *Anticipation of the Price of CAT Bond in the Management of the Process of Securitization of Catastrophe Insurance*, Econometrics 28, (ed.) P. Dittmann, Research Papers of Wrocław of Economics No. 91, Wydawnictwo Uniwersytetu Wrocławskiego, Wrocław 2010, s. 151-165.



Rys. 1. Ceny obligacji CAT w warunkach ryzyka odmowy zapłaty z moralnym hazardem ($\lambda, \sigma = (2,2)$, i $\phi = -3$)

Ceny CAT dla zmiennej intensywności straty i poziomu uruchomienia



Rys. 2. Ceny obligacji CAT o ryzyku odmowy wypłaty przy $A/L=1,1$ w warunkach moralnego hazardu oraz $\phi = -5$ i stałej zmienności σ_i

Powyższy rysunek odzwierciedla przypadek, gdy zmienność straty jest stała ($\sigma_i = 2$) i stosunek aktywów do zobowiązań $A/L = 1,1$. Z ilustracji przedstawionej na rysunku 3 można wywnioskować, iż w skrajnych przypadkach, gdy intensywność straty jest niska (na poziomie $\lambda = 0,5$), to niezależnie od poziomu uruchomienia wypłaty obligacji, składka za obligacje w warunkach ryzyka w uwarunkowaniach hazardu moralnego najbardziej odzwierciedli wypłatę w związku z poniesionymi stratami. Wzrost intensywności straty λ do poziomu 2 zdecydowanie zmniejsza wypłatę z obligacji.

Ceny obligacji CAT zilustrowane przestrzennie, z uwzględnieniem moralnego hazardu, przy odpowiednich wartościach parametrów (λ , σ_i) na rysunku 1 i rysunek 2 przedstawiający ceny obligacji CAT przy zmiennym poziomie uruchomienia i zmiennej intensywności straty i w warunkach istnienia hazardu moralnego wskazują na wyraźną zależność tych cen od ryzyka bazowego mierzonego stosunkiem A/L oraz od intensywności straty λ , co będzie wyraźnie widoczne w przypadku trzecim (rysunek 3). Rysunki te wskazują jednocześnie na niewielki wpływ na ceny CAT poziomu uruchomienia wypłat przy stałym ryzyku bazowym A/L . Istotne zmiany ceny wskazują natomiast, że moralny hazard jest ważny i powinien być brany pod uwagę przy wycenie obligacji CAT. Bantwal i Kunreuther³⁵ również wskazali, że moralny hazard może wyjaśniać pewne niepewności w szacowaniu składki obligacji CAT.

5.2. Numeryczne wyznaczenie cen obligacji CAT w warunkach ryzyka odmowy wypłaty z alternatywnym uwzględnieniem występowania hazardu moralnego

W tabeli 5 prezentowany jest wpływ ryzyka stopy bazowej na ceny obligacji CAT. Ponieważ różnica w wycenach CAT spowodowana elastycznością stopy procentowej jest mała, rozpatrywany jest tu jedynie przypadek, gdy $\phi = -3$ i uwaga skupiona jest na omówieniu wpływu ryzyka stopy bazowej. Tabela zawiera pełne wyniki jedynie dla poziomu uruchomienia $K = 100$ i 110 oraz współczynnika $A/L = 1,1$. Dla $A/L = 1,2$ i $1,3$ podane są tylko wartości obligacji CAT dla $(\lambda, \sigma_i, \sigma_{index}) = (2, 2, 2)$, a więc tylko dla największej intensywności strat i największej zmienności strat i indeksu straty.

³⁵ V.J. Bantwal, H.C. Kunreuther, op. cit.

Tabela 5

Ryzykowne ceny CAT bond względem bazowego ryzyka ($\lambda_r = -0,01$, $\Phi = -3$)

<i>Uruchomienia</i>	<i>K = 100</i>			<i>K = 110</i>		
$(\lambda, \sigma_i, \sigma_{index})$	$\rho_x = 0,5$	$\rho_x = 0,8$	$\rho_x = 1$	$\rho_x = 0,5$	$\rho_x = 0,5$	$\rho_x = 1$
	<i>A/L=1,1</i>					
(0,5, 0,5, 1)	0,95122	0,95122	0,95122	<i>0,95122</i>	<i>0,95122</i>	<i>0,95122</i>
(0,5, 1, 1)	0,94833	0,94854	0,94885	<i>0,94853</i>	<i>0,94867</i>	<i>0,94906</i>
(0,5, 2, 2)	0,89340	0,89808	0,97876	<i>0,89448</i>	<i>0,89859</i>	<i>0,90835</i>
(1, 0,5, 0,5)	0,95123	0,95123	0,95123	<i>0,95123</i>	<i>0,95123</i>	<i>0,95123</i>
(1, 1, 1)	0,94327	0,94383	0,94552	<i>0,94413</i>	<i>0,94467</i>	<i>0,94607</i>
(1, 2, 2)	0,83091	0,84144	0,85737	<i>0,83322</i>	<i>0,84387</i>	<i>0,86005</i>
(2, 0,5, 0,5)	0,95127	0,95127	0,95127	<i>0,95127</i>	<i>0,95127</i>	<i>0,95127</i>
(2, 1, 1)	0,92610	0,92876	0,93332	<i>0,92861</i>	<i>0,93075</i>	<i>0,93523</i>
(2, 2, 2)	0,71399	0,73115	0,76071	<i>0,71764</i>	<i>0,73464</i>	<i>0,76521</i>
$(\lambda, \sigma_i, \sigma_{index})$	<i>A/L=1,2</i>					
(2, 2, 2)	0,72460	0,74141	0,76736	<i>0,72943</i>	<i>0,74604</i>	<i>0,77496</i>
$(\lambda, \sigma_i, \sigma_{index})$	<i>A/L=1,3</i>					
(2, 2, 2)	0,73357	0,74992	0,77300	<i>0,73896</i>	<i>0,75536</i>	<i>0,78101</i>

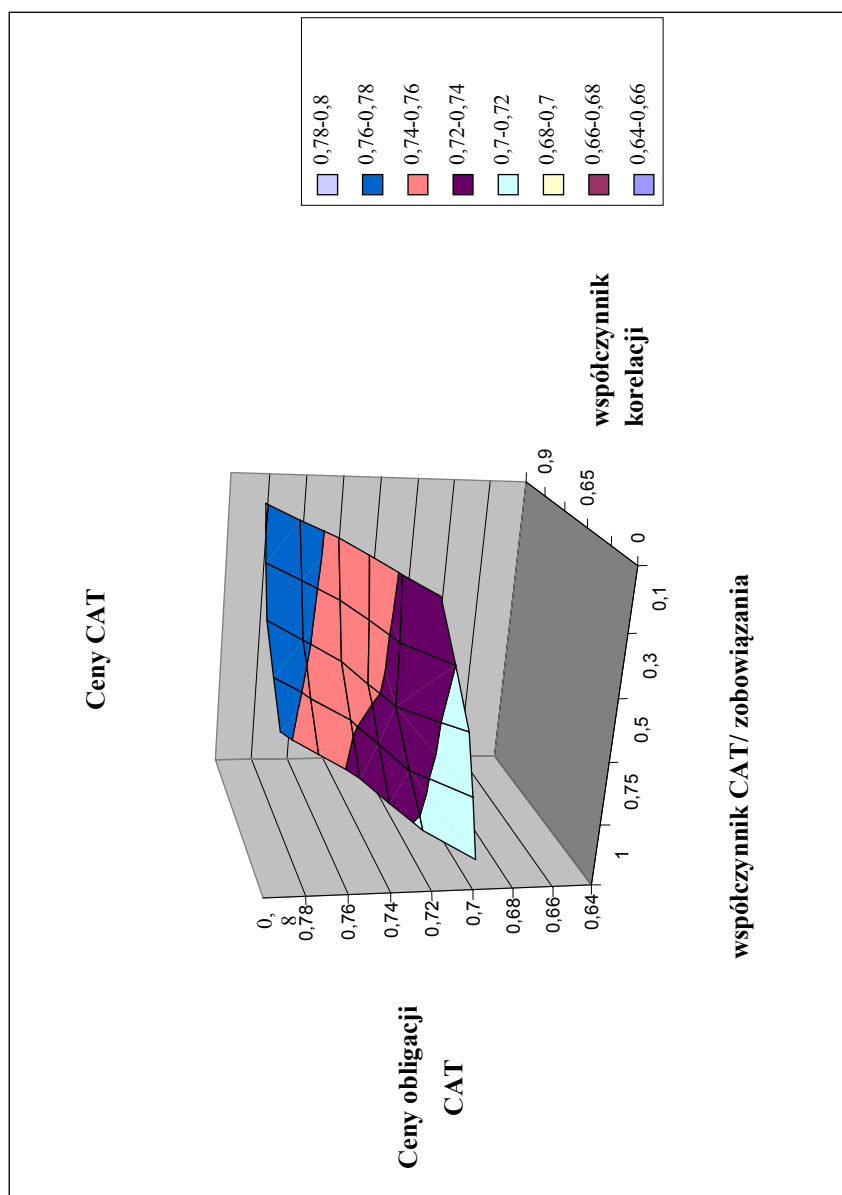
Źródło: Dane symulowane opracowane na podstawie: Jin-Ping Lee, Min-teh Yu, op. cit.

Należy zwrócić uwagę na fakt, iż przy niskim współczynniku korelacji pomiędzy poszczególną stratą a wskaźnikiem straty ρ_{index} , firma emitująca obligacje ma wysokie ryzyko bazowe. Kiedy $\rho_x = 1$, nie istnieje żadne ryzyko bazowe i oczekuje się, że ceny obligacji będą takie same jak ich odpowiednie wartości w tabeli 4. W przypadku $\rho_x = 0,8$ lub $0,5$ istnieje ryzyko bazowe i obserwujemy, że obniża to ceny obligacji CAT i że wielkość tego ryzyka wzrasta wraz z częstotliwością straty i zmiennością straty, co wpływa oczywiście na spadek wartości obligacji. Na przykład, w przypadku gdy $(\lambda, \sigma_i, \sigma_{index}) = (2, 2, 2)$, $K = 110$ i $A/L = 1,1$, cena obligacji CAT spadnie o 305 punktów bazowych, przy spadku wartości współczynnika ρ_x od 1 do 0,8. Natomiast spadek następuje o kolejne 170 punktów bazowych, gdy współczynnik ρ_x zmniejsza się z poziomu 0,8 do 0,5. Różnice ceny spowodowane przez ryzyko stopy bazowej są bardzo istotne w tym ustaleniu. Przy współczynniku A/L aktywów do zobowiązań równym 1,2 spadek cen obligacji CAT wynosi odpowiednio 290

i 166 punktów bazowych. Natomiast w najkorzystniejszym wypadku, gdy $A/L = 1,3$, spadki wartości obligacji CAT wynoszą odpowiednio: 256 i 164 punkty bazowe.

Należy także zauważyć, że wpływ ryzyka bazowego maleje wraz ze współczynnikami A/L i wzrasta wraz z poziomami uruchomienia, intensywnością straty i zmiennością straty. Rysunek 2 ilustruje tę sytuację. Można zauważyć, jaka jest relacja cen obligacji CAT i ryzyka bazowego (korelacji straty, ρ_x) oraz współczynników struktury zobowiązań (stosunku wartości emitowanych obligacji CAT do wartości ogólnej zobowiązań, a) dla przypadku $(\lambda, \sigma_i, \sigma_{index}) = (2, 2, 2)$, $K = 110$. Główny wniosek jaki tu się wyłania wskazuje, że cena obligacji CAT wzrasta wraz ze wzrostem stopy korelacji straty. Zatem korelacja straty zmniejsza składkę ryzyka bazowego przy rosnącej stopie korelacji straty. Wskazuje to także, że ubezpieczyciele z niską korelacją straty są podatni na znaczne dyskontowanie ich cen obligacji CAT.

Rysunek 3 wskazuje także, że cena obligacji CAT zmniejsza się wraz ze współczynnikiem struktury zobowiązań a przy rosnącej stopie, zwłaszcza przy niskich poziomach współczynnika A/L , co uwidacznia także rysunek 4, wskazując, że zobowiązania obligacji CAT zwiększają składkę odmowy wypłaty przy rosnącej stopie. A zatem, ubezpieczyciele mający więcej obligacji CAT obciążonej długiem, muszą płacić wyższą składkę za ich obligacje.



Rys. 3. Ceny obligacji CAT przy ryzyku odmowy wypłaty przy $A/L = 1,1$, $\phi = -3$, poziomie uruchomienia $K = 110$ oraz $\lambda_r = -0,01$

Podsumowanie

Rozpatrzony w artykule model wyceny obligacji CAT uwzględnia stochastyczne stopy procentowe i bardziej ogólne procesy straty katastrofalnej. Umożliwia zatem zmierzenie wpływów ryzyka odmowy wypłaty, moralnego hazardu i ryzyka stopy bazowej, jakie są związane z obligacjami CAT. W artykule³⁶ stwierdza się, że ceny obligacji CAT pozbawione ryzyka odmowy wypłaty i wyznaczone numerycznie są bardzo zbliżone do cen obliczonych przy zastosowaniu procedury aproksymacyjnej, z wyjątkiem sytuacji, kiedy zmienność straty jest wysoka. Wtedy aproksymowane ceny osiągają wyższe wartości.

W sekwencji trzech przypadków rozpatrzonych w niniejszym artykule szacowane są ceny obligacji CAT o ryzyku odmowy wypłaty i analizowana jest składka w warunkach ryzyka odmowy wypłaty, zmieniająca się wraz z wielkością roszczeń wywołanych danym zdarzeniem katastrofalnym, zmiennością straty, elastycznością stopy procentowej aktywów ubezpieczyciela, współczynnikiem początkowego kapitału i strukturą zobowiązań. Składka szacowana jest także ze względu na wpływ moralnego hazardu na ceny (wartość) obligacji CAT. Stwierdzono w tym kontekście, że moralny hazard w znacznym stopniu obniża wartość obligacji. Intensywność wpływu moralnego hazardu wzrasta wraz z intensywnością zdarzenia katastroficznego, zmiennością straty i ryzykiem stopy procentowej aktywów ubezpieczyciela, a maleje wraz z poziomem uruchomienia wypłat z obligacji oraz początkową wartością kapitału ubezpieczyciela.

W uzupełnieniu rozważań dotyczących zachowania się cen obligacji CAT jest uwzględnione ryzyko bazowe (systematyczne). Ryzyko bazowe znacząco zmniejsza ceny obligacji CAT i przyczynia się do zmniejszenia stopy tych cen. Natomiast oddziaływanie ryzyka bazowego wzrasta wraz z poziomem uruchomienia, intensywnością zdarzenia katastroficznego i zmiennością straty, a maleje z pozycją początkową kapitału.

Rozpatrzony w tym artykule model można postrzegać jako dość ogólny sposób oceny ryzyka dotrzymywania wypłaty zobowiązań. Wskazuje się tu przykładowe zastosowania tego modelu. Strukturalne ograniczenia w tym modelu odnoszą ceny obligacji do podstawowych charakterystyk aktywów, zobowiązań i stóp procentowych. Umożliwia to wycenę obligacji o unikatowych cechach przy zastosowaniu analizy numerycznej. Istotne jest przy tym, że model ten może być dość łatwo modyfikowany i zastosowany do analizowania oceny ryzyka dotrzymywania wypłaty innych zobowiązań, nie tylko dotyczących obligacji CAT, ale również papierów wartościowych związanych z ubezpieczeniem.

³⁶ *Ryzyko zobowiązań i działań...*, op. cit.

EFFECT OF BASE RATE AND THE RISK OF MORAL HAZARD ON THE PRICE OF CAT BONDS AND PAYMENT

Summary

The application of financial instruments of the capital market is aimed at managing the process of securitization of catastrophic risk. It is significant with regards to the results that stem from the loss caused by catastrophic phenomena. The article discusses the structure of CAT bonds and their pricing on the basis of the model of stochastic process of interest rate. CAT bonds are designed for financing the results of catastrophic events. They are similar to contingent claim CC, however, in reality, they are instruments of the financial market – catastrophic bonds.

The elaborated approach is illustrated by distribution of bond prices in the arrangement of the level of launching CAT bond remission and its variability. The article investigates the direction of research determined by the consideration of moral hazard and base risk (market risk) in the pricing.

In the sequence of cases considered in the article the prices of CAT bonds with the default risk are estimated. There is also the analysis of premium in the conditions of default risk that is changing along with the amount of claims caused by a given catastrophic event, variability of loss, flexibility of interest rate of the insurer assets, coefficient of initial capital and the structure of liabilities. The premium is estimated also with regard to the influence of moral hazard upon prices (value) of CAT bonds. In this context, it is stated that moral hazard significantly lowers the value of bonds. The intensity of influence of moral hazard increases together with the intensity of a catastrophic event, variability of loss and also the interest rate risk of the insurer assets and it is decreases along with the level of launching bond payment and the initial value of the insurer capital.

Włodzimierz Szkutnik

REGRESYJNE I DYNAMICZNO-STATYSTYCZNE MODELOWANIE ZMIENNOŚCI W PRAKTYCE OCENY RYZYKA INWESTYCJI KAPITAŁOWYCH I UBEZPIECZENIOWYCH

1. Modelowanie ekonometryczne i dynamiczno- -statystyczne

Wszystkie teorie makroekonomiczne wypracowane z pozycji modelowania ekonometrycznego w ostatnich dziesięcioleciach XX w. były oparte w znaczącej części na teorii Keynesa. Krytyka, jaka została zainicjowana w latach 70. XX w., wychodziła głównie z przesłanek krytyków keynesizmu, chociaż *explicite* nie formułowano tego stanowiska.

Mniej więcej od początku lat 80. ubiegłego wieku ekonometria wkroczyła w nową fazę rozwoju, gdy zaczęto modelować procesy finansowe, w tym wielkości notowane na giełdach. Wraz z modelem Engla (model ARCH – 1982 r.) i jego ucznia Bollersleva (model GARCH – 1986 r.) rozwinęły się w badaniach modele zmienności z autoregresyjną warunkową heteroskedastycznością.

Zmienność wielkości rynkowych jest bardzo ważna w zarządzaniu ryzykiem. Modelowanie zmienności dostarcza prostego ujęcia obliczania wartości narażonej na ryzyko VaR danej pozycji finansowej, która w różnych kontekstach aplikacyjnych pełni znaczącą rolę, a szczególnie w aplikacjach bankowych, co wynika z Umowy Kapitałowej (Bazylea). Modelowanie zmienności szeregu czasowego może prowadzić także do polepszenia skuteczności szacowania parametrów oraz poprawiać prognozowanie przedziałów zmienności.

W szczególnym, lecz ważnym przypadku, modele z warunkową heteroskedastycznością znajdują zastosowanie w badaniu zmienności rentowności aktywów. Szczególnie te modele będą przedmiotem dalszej analizy wynikającej w głównej mierze ze studiów literaturowych.

Oprócz zarysowanego charakteru i celu stosowania modeli typu GARCH opisującymi zmienność, ich zastosowanie jest ważne z punktu widzenia charakteru zmienności jako istotnego czynnika w handlu opcjami. W tym przypadku zmienność oznacza warunkową wariancję zasadniczej rentowności aktywów. Rozważając np. cenę *call option* (opcji kupna) typu europejskiego, posiadacz ma prawo, lecz nie obowiązek, kupna określonej liczby udziałów danej akcji zwykłej po ustalonej cenie w danym terminie. Ta ustalona cena nazywana jest *strike price* (cena uderzeniowa) i jest powszechnie w literaturze oznaczana przez K , natomiast „danym terminem” jest termin wygaśnięcia opcji. Ważne jest też znaczenie terminu „czasu trwania”, który jest okresem czasu pozostałym do terminu wygaśnięcia i oznaczanym przez t . Wiadomo natomiast, że gdy posiadacz może wykorzystać swoje prawo do kupna w terminie lub przed terminem wygaśnięcia polisy, to opcja jest typu amerykańskiego.

Wiadomo, że np. wycenę opcji typu europejskiego określa powszechnie już podawany w literaturze wzór Blacka-Scholesa. Wzór ten wykorzystamy w dalszej części artykułu, a poniżej bardziej szczegółowo ze względu na przyjętą tu konwencję, zostaną przedstawione modele zmienności.

1.1. Charakterystyka zmienności

Tak jak staraliśmy się wyrazić stanowisko dotyczące zmienności w artykule zamieszczonym w niniejszym zeszycie w kontekście kontrowersyjności dotyczących racjonalnych oczekiwań, jako pewien substytut niepewności lub ryzyka, tak w kontekście zmienności akcji, właściwie cen akcji notowanych na giełdzie lub szerzej traktując o tym, wartości aktywów, określimy teraz tę zmienność w aspekcie jej charakterystyki. Wyjściowym atrybutem zmienności jest to, że nie jest obserwowalna bezpośrednio, gdyż np. dzienna zmienność nie jest bezpośrednio zauważalna na podstawie zwrotów cen akcji. Praktycznie bowiem tylko jedna jest obserwacja w ciągu dnia. Inaczej byłoby, gdyby dane dotyczące akcji miały częstotliwość np. sekundową. Wtedy zmienność dzienną można by szacować. W praktyce, przy wymaganej poprawności oszacowań, należy dokładnie przebadать tę zmienność, tym bardziej, że zmienność dzienna akcji obejmuje zmienność o częstotliwości większej niż dzienna oraz zmienność między dniami handlowymi. Nieobserwowalność zmienności jest istotną przeszkodą w ocenie wyników prognozowania warunkowych modeli heteroskedastycznych.

Pewnym wyjściem z opisanej sytuacji, substytutem nieobserwowalności zmienności dziennej, jest wykorzystanie zasady, że ceny są regulowane przez model ekonometryczny, którego formą jest wzór Blacka-Scholesa, i wówczas

można cenę wprowadzić dla celów implikowania zmienności (tzw. domniemanej). Podejście takie ma jednak oczywiste wady ze względu na narzucenie określonego modelu wyznaczania cen, który z natury wymaga spełnienia pewnych założeń nie zawsze adekwatnych do obserwowanej w praktyce zmienności. Przekłada się to w oczywisty sposób na ocenę ryzyka wykrywania regularności takiej zmienności.

Przykład 1

Na podstawie zaobserwowanych cen opcji kupna typu europejskiego można, stosując wzór Blacka-Schoelsa, wyprowadzić warunkowe odchylenie standardowe σ_t . Wynikająca stąd wartość σ_t^2 nazywana implikowaną zmiennością jest wyprowadzona z lognormalnego rozkładu zwrotu szeregu cen akcji. W praktyce może się ona bardzo różnić od realnej zmienności, bowiem empiryczne analizy wskazują, że implikowana zmienność rentowności aktywów bywa większa niż ta wyprowadzona z modelu typu GARCH.

Zmienność, pomimo tego, że nie jest bezpośrednio obserwowana, ma pewne cechy, które są widoczne w przypadku analizy rentowności cen aktywów. Przede wszystkim obserwowane jest tzw. grupowanie zmienności, co znaczy, że zmienność w pewnych okresach może być bardzo wysoka, a w innych niska. Ponadto przebieg zmienności w czasie jest ciągły, co mimo obserwowanych anomalii powoduje, że skoki zmienności są zjawiskiem rzadkim. Wreszcie zmienność nie ma charakteru rozkładu α -stabilnego, co przekłada się na własność zmienności, że nie zmierza do nieskończoności i zmienia się w ramach pewnego przedziału. Ze statystycznego punktu widzenia oznacza to ponadto, że zmienność jest często stacjonarna. Zauważalną cechą charakteryzującą zmienność jest także to, że może reagować odmiennie na duże wzrosty i duże spadki ceny.

Wnioski

1. Powyższe własności mają ważne znaczenie w rozwoju modeli zmienności.
2. Pewne modele zmienności powstały tylko dla skorygowania nieadekwatności istniejących już modeli, po stwierdzeniu immanentnych cech charakteryzujących zmienność, np. model EGARCH skonstruowano dla uwzględnienia asymetrii zmienności wynikającej z dodatniej i ujemnej rentowności aktywów.

1.2. Struktura modelu

Niech r_t stanowi logarymiczną stopę zwrotu akcji w czasie t . Główną zasadą badań dotyczących zmienności jest to, że szereg ten nie wykazuje autokorelacji albo charakteryzuje go autokorelacja, lecz jest niższego rzędu i szereg nie jest zależny. Modele zmienności starają się uchwycić taką zależność w szeregach zwrotu.

Aby ukazać modele zmienności w odpowiedniej perspektywie, warto rozważyć średnią warunkową oraz warunkową wariancję logarymicznej stopy zwrotu r_t wyznaczoną przez F_{t-1} , tzn.

$$\mu_t = E(r_t | F_{t-1}), \quad \sigma_t^2 = Var(r_t | F_{t-1}) = E[(r_t - \mu_t)^2 | F_{t-1}] \quad (1)$$

gdzie F_{t-1} koduje zestaw informacji w czasie $t-1$.

Z badań empirycznych wynika, że F_{t-1} wyraża wszystkie liniowe funkcje minionych zwrotów. Zależność serii zwrotu akcji r_t jest słaba, o ile w ogóle występuje. Zatem równanie (1) dla μ_t powinno być nieskomplikowane i dlatego zakłada się, że r_t zmienia się tak, jak w prostym modelu szeregu czasowego, którym może być np. model ARMA (p, q) . Model ten ma znaną postać

$$r_t = \mu_t + a_t, \quad \mu_t = \Phi_0 + \sum_{i=1}^p \Phi_i \cdot a_{t-i} \quad (2)$$

dla r_t , gdzie p i q są nieujemnymi liczbami całkowitymi.

Model (2) ilustruje liczne możliwe zastosowania finansowe liniowych modeli szeregu czasowego. Struktura (p, q) modelu ARMA może zależeć od częstotliwości szeregu zwrotów. Na przykład, dzienne zwroty wskaźnika rynkowego często pokazują pewne drobne autokorelacje, lecz miesięczne zwroty wskaźnika mogą nie zawierać żadnej znaczącej autokorelacji. W modyfikowanym modelu możliwe jest włączenie pewnych zmiennych objaśniających do równania średniej warunkowej oraz wykorzystanie liniowego modelu regresji wraz z błędami szeregu czasowego, aby uchwycić zachowanie się średniej μ_t . Na przykład próbna zmienna może być zastosowana odnośnie do poniedziałków, aby zbadać wpływ weekendowy na dzienne stopy zwrotu akcji.

Łącząc równania (1) i (2), otrzymuje się model wariancji

$$\sigma_t^2 = Var(r_t | F_{t-1}) = Var(a_t | F_{t-1}) \quad (3)$$

Warunkowe modele heteroskedastyczne, których aspekt aplikacyjny ma istotne znaczenie, odnoszą się do rozwoju w czasie wariancji σ_t^2 . Sposób w jaki σ_t^2 jest modelowany, tzn. jego rozwój w czasie, jest podstawą do odróżniania od siebie modeli zmienności. Modele te mogą być zakwalifikowane do dwóch kategorii. Te z pierwszej kategorii urzeczywistniają kierowanie rozwojem zmienności σ_t^2 według dokładnej funkcji, natomiast w modelach zakwalifikowanych do drugiej kategorii zmienność jest uwzględniona przez równanie stochastyczne. I tak np. model GARCH należy do pierwszej kategorii, natomiast stochastyczny model zmienności – do drugiej kategorii.

W badaniach empirycznych zmienności dla uproszczenia zakłada się, że dany jest model średniej warunkowej. Jednak szacowane są zarówno równania średniej, jak i wariancji łącznie. Występujący w modelu element losowy a_t przyjęło się już nazywać szokiem (zwrotem odchyleniowym – ang. *shock*) lub zwrotem skorygowanym średnią (ang. *mean-corrected return*). Model dla μ_t w równaniu (3) jest nazywany równaniem średniej dla r_t , a model dla σ_t^2 równaniem zmienności dla r_t . Zatem modelowanie warunkowej heteroskedastyczności sprowadza się do sprowadzenia dynamicznego równania do modelu szeregu czasowego w celu kierowania w czasie rozwojem warunkowej zmienności odchylenia (ang. *of the shock*).

1.3. Stochastyczny model zmienności

Przykładem stochastycznego modelu zmienności jest znany z literatury¹ model losowego współczynnika autoregresji (ang. *random coefficient autoregressive* – RCA). Model ten jest najczęściej formułowany w celu wyjaśnienia zmienności w zakresie różnych przedmiotów badań. Sytuacja taka występuje także w analizie danych panelowych w ekonometrii oraz modelach hierarchicznych w statystyce. Model RCA klasyfikuje się jako warunkowy model heteroskedastyczny. Wydaje się jednak, co wynika z przeglądu literatury statystycznej, że właściwsze byłoby, także ze względów historycznych, zaklasyfikowanie modelu RCA do grupy modeli opisujących równanie warunkowej średniej procesu, w których parametry mogą się kształtować wraz z upływem czasu. Określa się taki sposób modelowania szeregu czasowego r_t jako naśladowujący model RCA(p). W przypadku modelu RCA warunkowa średnia i wariancja ma taką samą formę, jak w modelu CHARMA opisującym rozwój warunkowej wariancji.

¹ R.S. Tsay: *Analysis of Financial Time Series*, John Wiley & Sons, New York 2002.

Alternatywnym podejściem do opisu rozwoju zmienności finansowego szeregu czasowego jest wprowadzenie innowacji do równania warunkowej wariancji szoku losowego a_t . Wynikający stąd model jest znanym już dość powszechnie stochastycznym modelem zmienności SV (ang. *stochastic volatility*). Podobnie jak w modelach EGARCH, aby zapewnić dodatni znak wariancji, w modelach tych wprowadza się logarytm wariancji $\ln(\sigma_t^2)$, zamiast σ_t . Model ten definiowany jest w postaci

$$a_t = \sigma_t \cdot \varepsilon_t, \quad (1 - \alpha_1 \cdot B - \dots - \alpha_m \cdot B^m) \cdot \ln(\sigma_t^2) = \alpha_0 + v_t \quad (4)$$

gdzie ε_t są zmiennymi losowymi niezależnymi o identycznym rozkładzie $N(0,1)$, v_t są niezależnymi zmiennymi losowymi o rozkładzie $N(0, v_t^2)$ oraz niezależnymi od ε_t , α_0 jest stałą, a wszystkie pierwiastki wielomianu $1 - \sum_{i=1}^m \alpha_i \cdot B^i$ są co do wartości bezwzględnej większe niż 1.

Wprowadzenie innowacji v_t znacznie zwiększa elastyczność modelu w opisywaniu rozwoju wariancji σ_t^2 , także trudność w oszacowaniu parametru. Model SV może być szacowany metodą quasi-wiarygodności z zastosowaniem filtru Kalmana lub markowską metodą Monte Carlo (MCMC). Jacquier, Polon i Rossi porównują wyniki oszacowań tymi metodami².

1.4. Stochastyczne modele zmienności długiej pamięci

Dla praktycznych zastosowań ważna jest ta cecha modelu SV, która umożliwia zrekonstruowanie go w taki sposób, by uwzględnił długą pamięć zmienności. W tym celu wprowadza się do modelu tzw. różnice ułamkowe. Idea wynika z własności szeregu czasowego długookresowego, który jest taki wtedy, gdy jego charakter rozkładu współczynnika funkcji autokorelacji jest hiperboliczny, a nie jak się przyjmuje tradycyjnie, wykładniczy przy wzrastającym opóźnieniu. Rozbudowa w kierunku modeli o długiej pamięci w badaniach zmienności jest motywowana przez fakt, że funkcja autokorelacji szeregu kwadratowego lub szeregu wartości bezwzględnych rentowności aktywów często ma powolny rozkład, pomimo tego, że w szeregu zwrotu nie stwierdza się autokorelacji. Rozważają to Ding, Granger i Engle³. Przykładowa funkcja

² E. Jacquier, N. Polson, P. Rossi, *Bayesian Analysis of Stochastic Volatility Models*, „Journal of Business & Economic Statistics” 1994, Vol. 12, s. 371-417.

³ Z. Ding, C.W. Granger, R. Engle, *A Long Memory Property of Stock Return and a New Model*, „Journal of Empirical Finance” 1993, Vol. 1, s. 83-106.

autokorelacji dziennych absolutnych zwrotów akcji IBM oraz wskaźnika S&P 500 w latach 1926-1997 wskazuje na dodatnią wartość i umiarkowane wielkości przy powolnym rozkładzie.

Prosty model stochastyczny o długiej pamięci LMSV (ang. *long-memory stochastic volatility*) można wyrazić następująco:

$$a_t = \sigma_t \cdot \varepsilon_t, \quad \sigma_t = \sigma \cdot \exp(u_t / 2), \quad (1 - B)^d \cdot u_t = \eta_t \quad (5)$$

gdzie $\sigma > 0$, ε_t mają identyczny, niezależny rozkład $N(0,1)$, η_t mają niezależne rozkłady typu $N(0, \sigma_\eta^2)$ i niezależny od ε_t oraz $0 < d < 0,5$.

Cecha długiej pamięci wywodzi się z różnicy ułamkowej $(1 - B)^d$, która implikuje powolny rozkład funkcji autokorelacji przy hiperbolicznym, a nie wykładniczym współczynniku, kiedy następuje wzrost opóźnienia. Dla modelu (5) mamy

$$\ln(a_t^2) \equiv \mu + u_t + \varepsilon_t$$

Zatem szereg $\ln(a_t^2)$ logarytmów kwadratów szoków jest sumą gaussowskiego sygnału o długiej pamięci i niegaussowskiego białego szumu⁴. Oszacowanie stochastycznego modelu zmienności o długiej pamięci jest skomplikowane, lecz parametr różnicy ułamkowej d może być oszacowany metodą quasi-największej wiarygodności lub metodą regresji.

Przykładowo, stosując logarytmy szeregu dziennych kwadratowych stóp zwrotu dla firm o wskaźniku S&P 500, Bollerslev i Jubinski oraz Ray i Tsay⁵ ujawnili, że oszacowanie mediany dla d wynosi około 0,38. W celach aplikacyjnych zbadano także wspólne składniki długiej pamięci w dziennej zmienności akcji w grupach firm klasyfikowanych ze względu na różne cechy. Stwierdzono, że firmy w tym samym sektorze przemysłowym czy handlowym często mają więcej wspólnych składników długiej pamięci. Dotyczy to np. dużych narodowych banków USA i instytucji finansowych.

⁴ F.J. Breidt, N. Crato, P. de Lima, *On the Detection and Estimation of Long Memory in Stochastic Volatility*, „Journal of Econometrics” 1998, Vol. 83, s. 325-348.

⁵ T. Bollerslev, D. Jubinski, *Equality Trading Volume and Volatility: Latent Information Arrivals and Common Long-run Dependencies*, „Journal of Business & Economics Statistics” 1999, Vol. 17, s. 9-21; B.K. Ray, R.S. Tsay, *Long-Range Dependence in Daily Stock Volatilities*, „Journal of Business & Economics Statistics” 2000, Vol. 18, s. 254-262.

1.5. Podejście alternatywne – dynamiczność struktury dziennych stóp zwrotu i ich stochastyczność

Rozważania alternatywne względem modelowania zmienności o długiej pamięci były prowadzone także w kontekście danych o wysokiej częstotliwości w celu kalkulacji zwrotów o niskiej częstotliwości⁶. Podejście to ze względu na dostępność w obecnym okresie danych o dużej częstotliwości na powrót znalazło się w centrum uwagi badaczy szeregów stóp zwrotu. Dotyczy to zwłaszcza obszaru zagranicznych rynków wymiany walut⁷.

Dla celów prezentacji alternatywnego podejścia do szacowania zmienności przyjmijmy, że badana ma być miesięczna zmienność aktywów, dla których dostępne są dzienne zwroty. Niech r_t^m jest miesięczną logarymiczną stopą zwrotu aktywów w miesiącu t . Można założyć, że w miesiącu jest n dni handlowych, a dzienne logarymiczne stopy zwrotu aktywów w miesiącu to $\{r_{t,i}\}_{i=1,\dots,n}$. Z własności logarytmicznych wynika następująca równość:

$$r_t^m = \sum_{i=1}^n r_{t,i}$$

Zakładając istnienie warunkowej wariancji i kowariancji, mamy:

$$Var(r_t^m | F_{t-1}) = \sum_{i=1}^n Var(r_{t,i} | F_{t-1}) + 2 \cdot \sum_{i < j} Cov[(r_{t,i}, r_{t,j}) | F_{t-1}] \quad (6)$$

gdzie F_{t-1} zestaw kodowanej informacji do momentu $t-1$ (włącznie).

Z równania (6) przy założeniu, że szereg $\{r_{t,i}\}_{i=1,\dots,n}$ jest białym szumem, wynika własność

$$Var(r_t^m | F_{t-1}) = n \cdot Var(r_{t,1})$$

gdzie $Var(r_{t,1})$ szacowane jest z dziennych zwrotów $\{r_{t,i}\}_{i=1,\dots,n}$ przez

$$\hat{\sigma}_t^2 = \frac{\sum_{i=1}^n (r_{t,i} - \bar{r}_t)^2}{n-1}$$

⁶ K.R. Frensz, G.W. Schwert, R.F. Stambaugh, *Expected Stock Returns and Volatility*, „Journal of Financial Economics” 1987, Vol. 19, s. 3-29.

⁷ T.G. Andersen, T. Bollerslev, F.X. Diebold, P. Laby, *The Distribution of Exchange Rate Volatility*, Working paper, Economic Department, University of Pennsylvania 1999.

a $\bar{r}_t$ jest średnią z próby dziennych logarytmicznych stóp zwrotu w miesiącu t , tzn.

$$\bar{r}_t = \sum_{i=1}^n \frac{r_{t,i}}{n}$$

Zatem oszacowanie miesięcznej zmienności wyraża wzór

$$\hat{\sigma}_m^2 = \frac{n}{n-1} \cdot \sum_{i=1}^n (r_{t,i} - \bar{r}_t)^2 \quad (7)$$

Natomiast, jeśli $\{r_{t,i}\}_{i=1,\dots,n}$ naśladuje model autoregresji MA(1) z jedno-okresowym opóźnieniem, wówczas zmienność jest modelowana w równaniu

$$Var(r_t^m | F_{t-1}) = n \cdot Var(r_{t,1}) + 2 \cdot (n-1) \cdot Cov(r_{t,1}, r_{t,2})$$

Oszacowanie zmienności może być więc szacowane jak w wyrażeniu

$$\hat{\sigma}_m^2 = \frac{n}{n-1} \cdot \sum_{i=1}^n (r_{t,i} - \bar{r}_t)^2 + 2 \cdot \sum_{i=1}^{n-1} (r_{t,i} - \bar{r}_t) \cdot (r_{t,i+1} - \bar{r}_t) \quad (8)$$

Przedstawiony sposób szacowania zmienności jest prosty, lecz napotyka na pewne trudności w praktyce. Po pierwsze, model dla dziennych zwrotów $\{r_{t,i}\}_{i=1,\dots,n}$ nie jest znany. Komplikuje to oszacowanie wariancji w równaniu (6). Po drugie, w miesiącu jest około 21 dni handlowych, co znaczy, że próba ma małą licznosc. Zatem poprawność oszacowań wariancji i kowariancji w równaniu (6) może nie być właściwa. Poprawność ta zależy od dynamicznej struktury dziennych stóp zwrotu $\{r_{t,i}\}_{i=1,\dots,n}$ oraz od ich rozkładu. Należy tu zwrócić uwagę na fakt, że w przypadku gdy dzienne logarytmiczne stopy zwrotu mają wysokie ekscesy (kurtozy) i skorelowane serie wartości szeregu, wówczas oszacowania $\hat{\sigma}_m^2$ w równaniu (7) i (8) mogą nawet nie być stałe. Stwierdzają to w swych badaniach Bai, Russell i Tiao⁸. Wydaje się więc, że potrzebne są pewne modyfikacje zaprezentowanego alternatywnego podejścia w modelowaniu zmienności.

⁸ X. Bai, J.R. Russell, G.C. Tiao, *Beyond Metron's Utopia: Effectsof Dependence and Non-Normality on Variance Estimates Using High-Frequency Data*, Working paper, Graduate School of Business, University of Chicago, Chicago 2000.

Przykład 2

Rozważana jest miesięczna zmienność logarytmicznych stóp zwrotu wskaźnika S&P 500 od stycznia 1980 r. do grudnia 2009 r. Obliczona została zmienność za pomocą trzech metod. W pierwszej metodzie zastosowano dzienne logarytmiczne stopy zwrotu i równanie (7) (założono więc, że logarytmiczne stopy zwrotu są szeregiem opisanym białym szumem). W drugiej metodzie założono, że naśladują model MA(1), zastosowano więc równanie (7). Trzecia metoda modelowania zmienności została oparta na modelu GARCH (1,1) dla miesięcznych stóp zwrotu od stycznia 1962 r. do grudnia 2009 r. Dłuższy szereg był podyktowany celem uzyskania bardziej poprawnego oszacowania miesięcznej zmienności. Otrzymane oszacowanie modelu jest następujące:

$$r_t^m = 0,658 + a_t, \quad a_t = \sigma_t \cdot \varepsilon_t, \quad \sigma_t^2 = 3,349 + 0,086 \cdot a_{t-1}^2 + 0,735 \cdot \sigma_{t-1}^2,$$

gdzie ε_t jest standardowym szeregiem białego szumu. Czasowe oszacowanie miesięcznej zmienności wskazuje, że zmienności wyznaczone na podstawie dziennych stóp zwrotu są znacznie wyższe niż te oszacowane na podstawie miesięcznych stóp zwrotu czy te oszacowane na podstawie modelu GARCH (1,1). W szczególności, oszacowana zmienność dla października 1987 r. wynosiła 680, przy danych dziennych stopach zwrotu.

Uwaga 1

W równaniu (7), jeśli będzie założone, że średnia z próby $\hat{r}_t$ wynosi zero, wówczas

$$\hat{\sigma}_m^2 \approx \sum_{i=1}^n r_{t,i}^2$$

W tym przypadku skumulowana suma kwadratów dziennych logarytmicznych stóp zwrotu w miesiącu może być zastosowana jako oszacowanie miesięcznej zmienności.

1.6. Kierunki badań

Oprócz wskazanych wcześniej ograniczeń i zalet stochastycznych dynamicznych modeli zmienności można zastosować modele zmienności dla wskazania pewnych problemów wagi praktycznej.

1. Problem polega na stwierdzeniu, czy dzienna zmienność akcji jest niższa latem i jeśli tak, to jak bardzo. Twierdzące odpowiedzi na te dwa pytania, na podstawie miesięcznych logarytmicznych stóp zwrotu akcji badanej firmy, mają praktyczne implikacje w wycenie opcji na akcje.

2. Ze względu na to, że wskaźnik giełdowy jest powszechnie wykorzystywany na rynkach pochodnych, modelowanie jego zmienności jest inspirujące praktycznie. Zachodzi tu bowiem pytanie i pojawia się problem, czy przeszłe stopy zwrotu poszczególnych składowych wskaźnika giełdowego przyczyniają się do modelowania zmienności tego wskaźnika w obecności jego własnych stóp zwrotu. Dokładne badanie dotyczące tej kwestii jest inspirujące, jednak wykracza poza zakres poruszonych w tym opracowaniu zagadnień formalnych. Można jednak podjąć tę kwestię, używając przeszłych stóp zwrotu akcji badanej firmy jako zmiennych objaśniających.

1.7. Niepewność w oszacowaniu zmienności

Okazuje się, co jest z pozoru paradoksalne, że w oszacowaniu zmienności występuje niepewność. Następuje więc swoisty regres w koncepcji mającej ujawnić zmienność, czyli wielkości obarczonej ryzykiem. Niepewność w analizie zmienności jest często ignorowana. Generalnie związana jest ona z niestalością oszacowanej zmienności, co jest powszechnym elementem w modelowaniu wszelkich zjawisk, gdyż nigdy formalnie nie jest wiadome, na ile oszacowane modele są stałe, stabilne są ich parametry i oszacowane błędy oraz ewentualnie wyznaczone na ich podstawie prognozy.

W ocenie niestalości oszacowanej zmienności formalnie znajduje zastosowanie kurtoza. Okazuje się, że możliwe jest wyprowadzenie wzoru określającego eksces zmienności wyznaczanej dla modeli GARCH. To samo dotyczy innych wariantów tego modelu. W teorii wyznacza się kurtozę dla przypadków, gdy składnik szeregu resztowego ε_t ma lub nie ma rozkładu gaussowskiego. Zachowanie się wielkości szokowej a_t spełnia przy tym istotne znaczenie w aspekcie zachowania się końcówek rozkładu. W niektórych przypadkach rozkład jest gruboogonowy.

1.8. Praktyczne zagadnienia badawcze

Oprócz wskazanych już wcześniej obszarów badawczych pojawiają się różne inne ważne kwestie, których rozstrzygnięcie jest pomocne w zarządzaniu ryzykiem aktywów różnych firm. Do istotnych zagadnień należy badanie, czy opóźnione stopy zwrotu akcji firmy są użyteczne w modelowaniu zmienności wskaźnika. Z metodycznych zagadnień istotne są te dotyczące umiejętności dopasowania określonego typu modelu do istniejącego szeregu obserwacji. Wiele innych podyktowane jest wynikającymi z bieżących potrzeb analiz, których celem jest odpowiednie zmodyfikowanie analizy decyzyjnej do warunków zmiennego otoczenia firmy.

2. Analiza numeryczna zmienności wartości kontraktów terminowych z zastosowaniem modeli klasy ARCH/GARCH

Dla zilustrowania zmienności omówionej w pierwszej części głównie w kontekście modelowania SV przeprowadzono analizę stóp zwrotu z instrumentów finansowych, przyjmując za ich miarę zmienności wariancję lub odchylenie standardowe z badanej próby empirycznej. W wypadku zmienności instrumentów rynku kapitałowego są one składnikiem ryzyka rynkowego, które jest znaczące dla szeroko rozumianego ryzyka ekonomicznego. Modele uwzględniające miary zmienności zjawiska w zależności od czasu – są najczęściej stosowane w sferze ekonomii i finansów do analizy dyspersji badanych zjawisk.

W tej części artykułu będą przedstawione wyniki analizy numerycznej dla wewnątrzsesyjnej (ang. *intraday*) zmienności kontraktów terminowych notowanych na GPW w Warszawie przy założeniach teoretycznych modeli o warunkowej heteroskedastyczności. Dotyczą one kontraktów na USD i WIG20⁹. W tym przypadku możliwe było także zilustrowanie empiryczne zaprezentowanych modeli w celu porównania ich skuteczności w zakresie analizy wariancji składnika resztowego oraz błędu prognoz. Wszystkie prezentowane wyniki zostały otrzymane przy pomocy pakietu obliczeniowego GRET. Podstawą rozważań empirycznych były modele klasy ARCH i GARCH.

Zaprezentowanie możliwości wykorzystania tych modeli zostało wykonane na podstawie wewnątrzsesyjnych (czyli zaobserwowanych dla każdej zmiany wartości w trakcie trwania sesji giełdowej) rzeczywistych wartości kontraktów terminowych. Analizie zostały poddane szeregi czasowe dla kontraktów terminowych na USD (8167 obserwacji) oraz na WIG20 (40 769 obserwacji). Przedstawiono analizy wewnątrzsesyjnej zmienności wartości kontraktów terminowych na USD oraz WIG20 z datą wykonania 19/09/2008 dla sesji z okresu od 15/10/2007 do 19/09/2008.

Do analizy zmienności kontraktu terminowego na USD zostały zastosowane reszty z modelu ARMA (1,0). Model dał w wyniku następujące szacunki parametrów:

$$FUSD_t = 0,0472998 + 0,999793 FUSD_{t-1}$$

Do modelowania zmienności kontraktu terminowego na WIG20 zastosowano również model ARMA (1,0), którego postać jest następująca:

$$FW20_t = 0,665616 + 0,999742 FW20_{t-1}$$

⁹ W. Szkutnik, M. Pichura, *Analiza wewnątrzsesyjnej zmienności wartości kontraktów terminowych z zastosowaniem modeli klasy ARCH/GARCH*, Ekonometria, nr 24, Prognozowanie, UE, Wrocław 2009.

Okazało się, że szeregi czasowe wartości kontraktów terminowych w badanym okresie okazały się obarczone autokorelacją składników losowych. Stwierdzono to na podstawie kolerogramów, których wykresy nie będą tu prezentowane. Stwierdzono, że wartości funkcji ACF (autokorelacji) PACF (częstkowej autokorelacji) dla kontraktu na USD wyraźnie wskazują, że rząd opóźnień dla rozważanych tutaj modeli badania zmienności powinien być znacznie większy niż 1. Dla kontraktu na WIG20 można wysunąć podobny wniosek, lecz nie jest on już tak ewidentny. Z uwagi na fakt, że możliwości obliczeniowe pakietu GRETL dla modelu ARCH ograniczają się do zastosowania opóźnienia nie większego niż 2, tylko dla modelu GARCH zostały zastosowane wyższe rzędy opóźnień. W tym wypadku zwiększenie opóźnień często skutkuje jednak modelem o dużej współliniowości, co prowadzi do eliminacji takiego modelu w dalszej analizie.

Okazało się, że wyniki estymacji modeli ARCH i GARCH dla kontraktu terminowego na USD nie pozwalają na jednoznaczne określenie, jakie wartości opóźnień dają najlepsze dopasowanie do danych rzeczywistych. Dla większości sprawdzanych modeli współczynniki R^2 wynosiły ponad 0,999, jednak z kryterium Akaike (AIC) lub kryterium Bayesa wynikało, że do analizy zmienności kontraktu na USD należy zastosować modele ARCH(2) oraz GARCH(1,2).

Model ARCH(2) zmienności kontraktu na USD przyjął po oszacowaniu parametrów postać

$$\sigma_t^2 = 0,0661160 + 0,239576 a_{t-1}^2 + 0,116980 a_{t-2}^2$$

Estymacja modelu GARCH(1,1) dla zmienności kontraktu na WIG20 dała natomiast następujące wartości:

$$\sigma_t^2 = 0,00348166 + 0,299940 a_{t-1} + 0,0000000001 05 a_{t-1} + 0,700000 \sigma_{t-1}^2$$

Wartości kontraktu terminowego na USD uzyskane z powyższych modeli względem wartości rzeczywistych w badanym okresie wskazały, że w obydwu przypadkach można zaobserwować niewielkie odchylenia wyestymowanych wartości od wartości rzeczywistych. Oczywiście występują estymacje, które w znaczny sposób odbiegają od danych pochodzących z autentycznych notowań, jednak są one stosunkowo rzadkie.

Analiza zmienności kontraktu na WIG20 przy zastosowaniu zaproponowanych modeli wykazała, że zarówno dla modelu ARCH, jak i GARCH, użycie opóźnień o wartościach większych niż 1 pozwala na estymację parametrów przy wysokim poziomie istotności. Podobnie jak w wypadku badania zmienności kontraktu na USD, obydwa modele dla kontraktu na WIG20 wykazują bardzo wysokie wartości współczynnika R^2 (ponad 0,999). Jednakże najlepsze dopasowanie do danych obserwowanych w rzeczywistości, z ponownym zastosowaniem kryterium AIC oraz kryterium Bayesa, zostało osiągnięte dla modelu ARCH(2) oraz GARCH(1,1).

Model ARCH(2) zmienności kontraktu na WIG20 po oszacowaniu parametrów ma następującą postać:

$$\sigma_t^2 = 12,6062 + 0,0940014 a_{t-1}^2 + 0,134801 a_{t-2}^2$$

Estymacja modelu GARCH(1,1) dla zmienności kontraktu na WIG20 dała w efekcie model o parametrach

$$\sigma_t^2 = 0,0482125 + 0,0484425 a_{t-1} + 0,951558 \sigma_{t-1}^2$$

Wartości kontraktu terminowego na WIG20 uzyskane na podstawie obydwu badanych modeli w stosunku do wartości rzeczywistych zaobserwowanych w analizowanym okresie okazały, że wartości otrzymane z zastosowaniem modelu ARCH(2) mają nieznacznie mniejsze odchylenia od rzeczywistych w porównaniu do modelu GARCH(1,1). W obydwu przypadkach zaobserwowano jednak dobre dopasowanie wartości prognozowanych do zaobserwowanych w badanym okresie.

Prognozowanie zmienności obydwu kontraktów terminowych, zaprezentowane w artykule przy zastosowaniu zaproponowanych modeli, dało obiecujące wyniki. W obydwu przypadkach użyte modele dały bardzo dobre dopasowanie do danych rzeczywistych. Ponadto warto zauważyć, że dla obydwu kontraktów lepsze dopasowanie zostało uzyskane dla modeli klasy ARCH, stosując kryterium Akaike i kryterium Bayesa.

Podsumowanie

Omówione w tym artykule kwestie dotyczące modelowania zmienności cen aktywów w postaci stochastyczno-dynamicznej wykazują zmienność stóp zwrotu aktywów i umożliwiają opracowanie prognozy odpowiednich wskaźników oceniających rentowność aktywów. Wybrane modele, takie jak znany z literatury model SV, miały na celu zilustrowanie idei analizy rentowności aktywów i jednocześnie wskazania ich uogólniających własności. Rozpatrywane zagadnienia są bardzo szeroko analizowane za pomocą licznych narzędzi analitycznych. Wielość modeli, które mają zastosowanie w badaniu zmienności aktywów, wynika z doniosłości badanych zagadnień. Dla rozwoju gospodarczego, stabilnego i obciążonego ograniczonym ryzykiem, zmienność implikująca ryzyko, a także niepewność, chociaż występuje odwrotny pierwotny wpływ tych wielkości, ma niezaprzeczalne znaczenie. Przenikanie się rynku finansowego i rynków ekonomicznych wymusza ciągłe studia analityczne nad badaniem różnorodnych zależności, gdyż analizy o charakterze opisu sytuacji prze-

szłej i bieżącej, podparte nawet bardzo subtelnymi teoriami ekonomicznymi stają się niewystarczające dla wyjaśnienia zjawisk finansowych i z płaszczyzny gospodarczej. Nie jest to zależne od tego, czy określone lobby naukowe propaguje określoną ideologię metodologiczną.

Zaprezentowane wyniki empirycznej analizy świadczą natomiast, że zastosowanie modeli zarówno klasy ARCH, jak i GARCH jest uzasadnione. W szczególności dotyczy to rozważanej tu zmienności cen instrumentów finansowych o dużej częstotliwości zmian. Uwzględniając wewnątrzsesyjne notowania kontraktów terminowych, stwierdzono, że wyniki mogą być obiecujące. Wskazuje na to wręcz wzorcowe dopasowanie wartości modelowych do rzeczywistych wartości notowań. W analizowanym okresie zaobserwowano niewielkie odchylenie od obserwowanych wartości. Dla danych empirycznych będących podstawą estymacji zmienności okazało się, że lepsze dopasowanie uzyskano na podstawie modelu ARCH. Nie jest to oczywiście ogólna prawidłowość. Analizując zjawiska ekonomiczne, szczególnie z rynku finansowego, należy testować modele na większej liczbie instrumentów kapitałowych, w szerszym i bardziej zróżnicowanym horyzoncie czasowym. Należy zwrócić też uwagę, że analiza dotycząca zmienności, której dokonano w artykule, umożliwia prognozowanie wariancji, co daje podstawy do predykcji i oceny przyszłego kształtowania kontraktów terminowych.

REGRESSION AND STATISTICAL MODELING OF DYNAMIC AND VARIATION IN PRACTICE, RISK ASSESSMENT, CAPITAL INVESTMENTS AND INSURANCE

Summary

Discussed in the article typical stochastic-dynamic models are formulated taking into account the volatility of asset returns and allow for longer term development of appropriate indicators evaluating forecasting profitability of assets. Some models, such as known from the literature SV model, aimed to illustrate the idea of the profitability analysis of assets and simultaneously indicate its generalization property. Issues are dealt with very extensively analyzed using several analytical tools. The multiplicity of models that are applicable in the study of variation of assets due to the significance of the issues studied. For economic development, stable and saddled with limited risk, the undeniable importance of risk and volatility implying the uncertainty, though in the real economy is the primary impact of these opposite magnitude. Interpenetration of financial markets and economic markets requires the development of new analytical techniques for the study of various quantities characterizing the relationship between these markets. In the second part of the article presents the analysis within a session (with English

intraday) volatility of the value of futures contracts on WIG20 USD and the date of exercise session 09.19.2008 dating from 15/10/2007 to 19/09/2008. Illustration of empirical models presented, allowed to formulate the conclusion that they can be a useful tool in the study of variation in the financial markets, in particular, phenomena characterized by high frequency changes.

Bartosz Lawędziak

OCENA SEKURYTYZACJI JAKO PROCESU ODDZIAŁYWUJĄCEGO NA STABILNOŚĆ FINANSOWĄ SYSTEMU BANKOWEGO

Wprowadzenie

Sekurytyzacja w ostatnich latach odgrywała istotną rolę w rozpraszaniu ryzyka kredytowego, tym samym podkreślając elastyczność systemu na ewentualne zaniedbania ze strony pożyczkobiorców. Ostatni kryzys kredytowy zrewidował to spojrzenie, głównie ze względu na fakt przedostania się zagrożonych pożyczek przez system finansowy oraz udostępnienie ich niczego niepodejrzewającym inwestorom/kredytobiorcom.

Artykuł powstał na podstawie studiów literaturowych z zakresu sekurytyzacji, kryzysu finansowego i podstaw księgowości. W pierwszej części została przedstawiona geneza sekurytyzacji i zależności z sprawozdaniami finansowymi jako całość konstrukcji szkieletu finansowego¹. Następnie siatka powiązań finansowych została nałożona na model struktury finansowej² bazujący na wartości ryzyka, co umożliwia zobrazowanie boomu pożyczkowego napędzanego spadkiem mierzonego ryzyka. Teoretyczny charakter artykułu został podyktowany brakiem rozwiązań sekurytyzacyjnych w sektorze bankowym, w związku z tym przedstawiono tylko pewne scenariusze działania. Na koniec zostały przedstawione wnioski dotyczące sekurytyzacji i jej udziału w stabilności finansowej podmiotów finansowych.

1. Podstawy sekurytyzacji

Od lat 80. w USA przedsiębiorstwa sponsorowane przez rząd (GSE), takie jak Fannie Mae lub Freddie Mac stały się największymi posiadaczami hipotek na nieruchomości oraz zauważalnie rosła ilość instrumentów sekuryty-

¹ Pojęcie wymiennie stosowane z zasadami rachunkowości.

² Rozumiana jako stosunek aktywów do poziomu zadłużenia.

zacyjnych w przedsiębiorstwach emitujących zabezpieczone aktywa (ABS). Pozabankowi emitenci ABS są posiadaczami hipotek, które nie spełniały wymogów GSE – hipoteki małe, jak i wielkie (tzw. słoniowe), wykraczające poza standardowe progi GSE.

Od chwili rozpoczęcia kryzysu kredytowego 2007/2008 pojawił się pogląd, który podkreślał wielowarstwowe problemy agencyjne obejmujące każdy poziom procesu sekurytyzacji, poczynając od źródła uzyskania pożyczki do jej sprzedaży³. Przede wszystkim nie ustalono granicy pomiędzy sprzedawaniem złych pożyczek w dół łańcucha i emitowaniem zobowiązań opartych na złych pożyczkach. Poprzez sprzedanie zagrożonej pożyczki pozbywano się jej ze swojego sprawozdania finansowego, natomiast emitowanie zobowiązań przeciw złym kredytom nie wyklucza ich z bilansu finansowego. Pośrednicy finansowi zaczęli w niekontrolowany sposób wykorzystywać instrumenty emitowane przez SPV, stając się ważnym fragmentem struktury finansowej narażonym na utratę swoich udziałów. Znajduje to potwierdzenie w rzeczywistości, bowiem z ok. 1,4 biliona dolarów oprocentowanych kredytów hipotecznych, połowa potencjalnych strat jest spowodowana przez amerykańskie struktury finansowe, takie jak banki komercyjne, banki inwestycyjne i zabezpieczone fundusze finansowe, a uwzględniając zagraniczne struktury finansowe – bilans rośnie do 2/3⁴.

2. Zasady rachunkowości

Pośrednicy finansowi odgrywają podwójną rolę, zarówno pożyczkodawcy, jak i pożyczkobiorcy. W celu uwzględnienia roszczeń oraz obligacji opisany zostanie teraz odpowiedni szkielet księgowy⁵. Przyjmując, że jest $n+1$ udziałowców w systemie finansowym, gdzie n z nich jest strukturami finansowymi (nazwijmy je dla uproszczenia bankami) oraz jeden nie jest strukturą finansową (oznaczony jako $n+1$), który gromadzi sprawozdania finansowe instytucji niebędących strukturami, jak np. towarzystwa ubezpieczeniowe lub fundusze inwestycyjne oraz domowi inwestorzy i zagraniczne banki centralne. Jest również sektor użytkowników końcowych, którymi są najwięksi pożyczkobiorcy. Dla naszych celów pożyczkobiorcami mogą być ludzie, którzy biorą kredyty hipoteczne na zakup domów.

³ A. Ashcraft; T. Schuermann, *Understanding the Securitization of Subprime Mortgage Credit*, Staff Report 2008, No. 318, Federal Reserve Bank of New York, http://www.newyorkfed.org/research/staff_reports/sr318.pdf

⁴ D. Greenlaw, J. Hatzius, A. Kashyap, H.S. Shin, *Leveraged Losses: Lessons from the Mortgage Market Meltdown*, Report of the US Monetary Monetary Form, No. 2, <http://www.chicagogs.edu/usmpf/docs/usmpf2008confdraft.pdf>, 2008

⁵ Na podstawie: H.S. Shin, *Securitisaton and financial stability*, „The Economic Journal” 2009, Vol. 119.

Oznaczmy przez $\bar{y}_i$ wartość nominalną roszczeń trzymanyh przez bank i względem użytkowników końcowych – pożyczkobiorców. Tak jak pożyczki użytkowników końcowych, istnieją również zobowiązania pomiędzy instytucjami finansowymi. Zobowiązanie jednej części w systemie będzie aktywem innej. Oznaczmy przez $\bar{x}_i$ wartość nominalną obligacji banku i oraz przez π_{ij} udział obligacji i -tego banku trzymanyh przez bank j . Jako $\bar{e}_i$ będziemy rozumieć wartość hipotetyczną udziałów banku i , co w bilansie finansowym banku można określić jako:

$$\bar{y}_i + \sum_{j=1}^n \bar{x}_i \pi_{ji} = \bar{x}_i + \bar{e}_i \quad (1)$$

Lewa strona równania (1) jest hipotetyczną wartością całkowitych aktywów banku i , składających się z pożyczek udzielonych użytkownikom końcowym $\bar{y}_i$ oraz roszczeniom przetrzymywanym wobec innych instytucji finansowych (banków) $\sum_{j=1}^n \bar{x}_i \pi_{ji}$. Prawa strona pokazuje całkowite zobowiązanie banku i wyrażone w wartości hipotetycznej oraz zawiera całkowita obiecaną spłatę $\bar{x}_i$ przez bank i plus hipotetyczne udziały $\bar{e}_i$. Zablokowane roszczenia oraz obligacje zostały przedstawione w tabeli 1, gdzie $\bar{x}_{ij}$ oznacza hipotetyczną wartość obligacji banku i do banku j .

Tabela 1

Macierz roszczeń i obligacji międzybankowych

	bank 1	bank 2	...	bank n	zewnętrzne	dług
bank 1	0	$\bar{x}_{12}$	...	$\bar{x}_{1n}$	$\bar{x}_{1,n+1}$	$\bar{x}_1$
bank 2	$\bar{x}_{21}$	0	...	$\bar{x}_{2n}$	$\bar{x}_{2,n+1}$	$\bar{x}_2$
...	...	...	...	...	...	...
bank n	$\bar{x}_{n1}$	$\bar{x}_{n2}$	...	0	$\bar{x}_{n,n+1}$	$\bar{x}_n$
pożyczki użytkowników końcowych	$\bar{y}_1$	$\bar{y}_2$	...	$\bar{y}_n$		
całkowite aktywa	$\bar{a}_1$	$\bar{a}_2$	...	$\bar{a}_n$		

Źródło: Na podstawie H.S. Shin, *Securitisation and Financial Stability*, „The Economic Journal” 2009, Vol. 119.

Sumowanie i -tego wiersza macierzy daje całkowite zobowiązanie banku i , a suma wpisów w i -tej kolumnie macierzy daje całkowitą wartość hipotetyczną aktywów banku i ; całkowita hipotetyczna wartość aktywów banku i jest oznaczona jako $\bar{a}_i$.

2.1. Ryzyko kredytowe

Przyjmujemy, że pożyczka zostaje wzięta w dacie „0”, a oddana ma zostać w dacie „1”. Pożyczki udzielone użytkownikom końcowym są obarczone ryzykiem, z którym banki muszą się liczyć. Uwzględnia ono jeden model współczynnikowy⁶, który jest szeroko stosowany i został zaadaptowany jako podstawa regulacji kapitałowych Basel II⁷. Zgodnie z modelem Vasicka pożyczkobiorca j , pożyczający od banku i spłaca dług, kiedy zmienna Z_{ij} jest nieujemna i zdefiniowana jako:

$$Z_{ij} = -\Phi^{-1}(p_i) + \sqrt{\rho}Y + \sqrt{1-\rho}X_{ij} \quad (2)$$

gdzie Φ jest dystrybucją rozkładu normalnego, Y i $\{X_{ij}\}$ są obustronnie zależnymi, normalnymi zmiennymi losowymi, ρ i p_i są stałymi. Y jest interpretowany jako ogólny czynnik ryzyka, X_{ij} jako specyficzny czynnik ryzyka. Prawdopodobieństwo zaniebdania p_i ze strony pożyczkobiorcy j wobec banku i jest określone następująco:

$$\Pr(Z_{ij} < 0) = \Pr[\sqrt{\rho}Y + \sqrt{1-\rho}X_{ij} < \Phi^{-1}(p_i)] = \Phi[\Phi^{-1}(p_i)] = p_i$$

Zależnie od czynnika Y , zaniebdania są niezależne wśród pożyczkobiorców i parametr ρ daje powiązanie *ex ante* w zaniebdaniach pomiędzy każdymi dwiema pożyczkami udzielonymi przez bank i . Zakładamy dalej, że portfel papierów wartościowych banku i zawiera N pożyczek dla użytkowników końcowych wyrażonych jako wartość nominalna y_i/N . Uznając, że N jest odpowiednio duże, przyjmujemy, że w portfelu papierów wartościowych pojawia się dużo małych pożyczek, których możliwe zaniebdania przez pożyczkobiorców zależą warunkowo od wykonania zmiennej Y . Na podstawie prawa wielkich liczb spłata ω_i w puli pożyczek wartości nominalnej $\bar{y}_i$ staje się funkcją deterministyczną zmiennej Y :

$$\omega_i \equiv \bar{y}_i \Pr(Z_{ij} \geq 0 | Y) = \bar{y}_i \Pr[Y\sqrt{\rho} + X_{ij}\sqrt{1-\rho} \geq \Phi^{-1}(p_i)] = \bar{y}_i \Phi\left[\frac{Y\sqrt{\rho} - \Phi^{-1}(p_i)}{\sqrt{1-\rho}}\right],$$

⁶ O. Vasicek, *The Distribution of Loan Portfolio Value*, 2002, http://www.moodyskmv.com/conf04/pdf/papers/dist_loan_port_val.pdf

⁷ A. Alizalde, R. Repullo, *Economic and Regulatory Capital in Banking: What is the Difference?*, Working paper, CEMFI, 2006.

a funkcja dystrybuanty spłaty dla banku i wyraża się wzorem:

$$F_i(z) = \Pr[\omega_i(Y) \leq Z] = \Pr\left[Y \leq \omega_i^{-1}(z) = \Phi\left[\frac{\Phi^{-1}(p_i) + \sqrt{1-\rho}\Phi^{-1}(z/\bar{y}_i)}{\sqrt{\rho}}\right]\right] \quad (3)$$

Zmiana p_i (prawdopodobieństwo zaniedbania na pojedynczej pożyczce wydanej przez bank i) sugeruje zmianę stochastyczną pierwszego stopnia w gęstości spłat. Kiedy p_i jest stałe, główna spłata pozostaje niezmienną, ale zmiana parametru ρ przy niezmiennym parametrze p_i sugeruje dominację stochastyczną drugiego rzędu w gęstości spłat. Wzrost parametru ρ jest połączony z marżą ochronną w gęstości spłat, czyniąc pulę pożyczek bardziej ryzykowną.

2.2. Realna wartość zadłużenia

Wiadome jest, że realna wartość spłat pożyczek przez użytkowników końcowych determinuje realną wartość roszczeń pomiędzy bankami, ponieważ zdolność jednego banku do spełnienia swoich obietnic jest zależna od środków, jakie ma zgromadzone na obligacjach. W tej części zostanie opisana realna wartość w dacie „1”. Przyjmijmy następujące oznaczenia: $\hat{y}_i$ – realna wartość spłat pożyczek banku i przez pożyczkobiorców, a $\hat{x}_i$ – realna wartość spłat pożyczek przez bank i . Zakładając, że każdy dług jest tak samo wymagalny, więc $\hat{x}_i < \bar{x}_i$, to wtedy bank j otrzyma π_{ij} akcji z $\hat{x}_j$. Wierzyciele otrzymują pełną wartość aktywów banku, jeśli realna wartość aktywów nie przekracza wartości nominalnej, stąd realną wartość długu opisuje zależność:

$$\begin{aligned} \hat{x}_1 &= \min(a_1(\hat{x}), \bar{x}_1) \\ \hat{x}_2 &= \min(a_2(\hat{x}), \bar{x}_2) \\ &\vdots \\ \hat{x}_n &= \min(a_n(\hat{x}), \bar{x}_n) \end{aligned} \quad (4)$$

gdzie: $\hat{x} = (\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n)$ jest profilem realnej wartości zadłużenia. Istnieje nie-malejąca funkcja $F(\cdot)$, która określa realną wartość aktywów do realnej wartości aktywów. Alokacja *ex post* jest stałym punktem na wykresie $F(\cdot)$. W warunkach średniej regularności istnieje jeden unikalny stały punkt dla prostej ekspozycji. Ponadto, biorąc unikalny stały punkt na funkcji $F(\cdot)$, realna wartość długu i -tego banku może zostać zapisana jako funkcja realnych spłat

pożyczek przez odbiorców końcowych $\hat{y} = (\hat{y}_1, \hat{y}_2, \dots, \hat{y}_n)^8$. Ponieważ realne wartości funkcji $\{\hat{y}_i\}$ są deterministycznymi funkcjami Y , można zapisać realną wartość aktywów banku i jako deterministyczną funkcję czynnika Y , a stąd:

$$\hat{a}_i(Y) = \hat{y}_i(Y) + \sum_j \pi_{ji} \hat{x}_j[Y(Y)] \quad (5)$$

Podsumowując, unikalny stały punkt na wykresie $F(\cdot)$ wzrasta w realnych spłatach $\{\hat{y}_i\}$, zatem $\hat{x}_j(\hat{y})$ jest rosnącą funkcją $\hat{y}$, co oznacza, że dla każdego banku i , realna wartość jego aktywów $\hat{a}_i$ jest dobrze zdefiniowaną rosnącą funkcją czynnika Y .

2.3. Wartość rynkowa

Wartość rynkowa jest definiowana jako oczekiwana wartość (od daty „0”) możliwych wartości realnych w dacie „1”. Do opisanie wartości pożyczek udzielonych przez bank i odbiorcom końcowym używamy y_i . Podobnie x_i jest wartością rynkową długów banku i , wziętą jako oczekiwaną wartość realnej wartości długu $\hat{x}_j$ w dacie „1”. Całkowita wartość rynkowa aktywów banku i może być wtedy zapisana jako:

$$a_i = y_i + \sum_j x_j \pi_{ji} \quad (6)$$

Bilans finansowy dla wartości rynkowej banku i dana jest wzorem:

$$y_i + \sum_j x_j \pi_{ji} = e_i + x_i \quad (7)$$

Lewa strona równania jest wartością rynkową aktywów, a prawa strona równia jest wartością rynkową zobowiązań w bilansie finansowym, gdzie e_i oznacza wartość rynkową udziałów banku i . Macierz roszczeń i obligacji pomiędzy bankami może być wtedy zapisana w wartościach rynkowych, tak jak pokazano w tabeli 2. i -ty wiersz macierzy może być sumowany w celu uzyskania wartości rynkowej długu banku i , a i -ta kolumna macierzy może być sumowana w celu uzyskania całkowitej wartości rynkowej aktywów banku i .

⁸ L. Eisenberg, T.H. Noe, *Systemic Risk in Financial Systems*, „Management Science” 2001, Vol. 47.

Tabela 2

Macierz roszczeń i obligacji pomiędzy bankami

	bank 1	bank 2	...	bank n	zewnętrzne	dług
bank 1	0	x_{12}	...	x_{1n}	$x_{1,n+1}$	x_1
bank 2	x_{21}	0	...	x_{2n}	$x_{2,n+1}$	x_2
...	...	...	...	...	...	...
bank n	x_{n1}	x_{n2}	...	0	$x_{n,n+1}$	x_n
pożyczki użytkowników końcowych	y_1	y_2	...	y_n		
całkowite aktywa	a_1	a_2	...	a_n		

Źródło: Na podstawie H.S. Shin, *Securitisaton and Financial Stability*, „The Economic Journal” 2009, Vol. 119, s. 317.

Z tożsamości (7) wektor wartości długu wśród banków można zapisać zgodnie ze wzorem (8), gdzie Π jest macierzą $n \times n$, gdzie (i, j) jest wartością jest π_{ij} .

$$[x_1, \dots, x_n] = [x_1, \dots, x_n] \Pi + [y_1, \dots, y_n] - [e_1, \dots, e_n] \quad (8)$$

lub bardziej zwięźle jako:

$$x = x\Pi + y - e \quad (9)$$

Równanie (9) pokazuje rekurencyjną naturę długu w systemie finansowym. Wartość długu każdego banku jest rosnąca w wartości długu innych banków. Rozwiązując równanie ze względu na y , otrzymujemy:

$$y = e + x(I - \Pi)$$

Strukturę finansową λ_i banku i można zapisać zgodnie z wzorem (10) jako współczynnik rynkowej wartości aktywów, do rynkowej wartości jego udziałów

$$\lambda_i \equiv \frac{a_i}{e_i} \quad (10)$$

Ponieważ $x_i / e_i = \lambda_i - 1$, $x = e(\Lambda - I)$, gdzie Λ jest macierzą diagonalną, której i -ty diagonalny wpis jest λ_i , stąd

$$y = e + e(\Lambda - I)(I - \Pi) \quad (11)$$

Dlatego profil całkowitych pożyczek udzielonych przez n banków do końcowych pożyczkobiorców zależy od interakcji w trzech podstawowych cechach systemu bankowego – dystrybucji udziałów e w systemie bankowym, profilu struktury finansowej Λ i struktury systemu finansowego zdefiniowanej przez Π , jak można by tego oczekiwać. Bardziej subtelna jest rola systemu finansowego dana przez macierz Π . Definiując wektor z jako:

$$z \equiv (I - \Pi)u \quad (12)$$

gdzie $u \equiv \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}$, stąd $z_i = 1 - \sum_{j=1}^n \pi_{ij}$. Innymi słowy z_i jest proporcją długu i -tego

banku, będącego w rękach zewnętrznych pożyczkodawców, sektora $n+1$. Wtedy całkowita ilość pożyczek udzielonych odbiorcom końcowym $\sum_i y_i$ może być wyrażona poprzez pomnożenie wzoru (11) przez u , skąd wynika

$$\sum_{i=1}^n y_i = \sum_{i=1}^n e_i + \sum_{i=1}^n e_i z_i (\lambda_i - 1) \quad (13)$$

Równanie (13) jest tożsamością bilansu finansowego dla sektora finansowego jako całości, gdzie wszystkie roszczenia i obligacje między bankami zostały sprowadzone do poziomu netto. Lewa strona równania to całkowita ilość pożyczek wziętych przez odbiorców końcowych. Pierwszy wyraz po prawej stronie równania (13) jest całkowitym udziałem systemu bankowego, a drugi wyraz po prawej stronie jest całkowitym finansowaniem zapewnianym przez sektor pozabankowy (co można zapisać jako $\sum_{i=1}^n x_i z_i$). Ostatecznie zapewnienie kredytów odbiorcom końcowym wynika zarówno z udziałów banku, jak i funduszy spoza sektora bankowego.

2.4. Struktura systemu finansowego

Przyjęty stopień struktury systemu finansowego jako całości jest zgodny z szeroką paletą poziomów struktur finansowych indywidualnych banków, co jest prawdziwe zarówno dla wartości nominalnych, jak i dla wartości rynkowych. System finansowy w wartościach nominalnych może być reprezentowany jako tablica $(\bar{e}, \bar{y}, \bar{x}, \Pi)$, która spełnia tożsamość finansową:

$$\bar{x} = \bar{x}\Pi + \bar{y} - e \quad (14)$$

Następnie dla dodatniej stałej ϕ można skonstruować system finansowy, którego sumaryczne udziały, pożyczki i struktura finansowa zostają niezmienione, ale gdzie stosunek długu do udziałów wszystkich indywidualnych banków jest ϕ razy większy. Zwłaszcza rozważając system finansów $(\bar{e}', \bar{y}', \bar{x}', \Pi')$, gdzie $\bar{e}' = \bar{e}$, $\bar{x}' = \Phi \bar{x}$ i Π' jest dowolną macierzą międzybankowych zobowiązań, której suma i -tego wiersza $1 - z_i / \Phi$, ostatecznie więc $\bar{y}'$ jest zdefiniowane jako

$$\bar{y}' = \bar{e}' + \bar{x}'(I - \Pi') \quad (15)$$

Uwzględniając powyższe, sumaryczna ilość pożyczek jest dana jako:

$$\sum_{i=1}^n \bar{y}'_i = \bar{e}'u + \bar{x}'(I - \Pi')u = \sum_{i=1}^n \bar{e}'_i + \sum_{i=1}^n \bar{x}'_i \frac{z_i}{\phi} = \sum_{i=1}^n \bar{e}'_i + \sum_{i=1}^n \bar{x}'_i z_i = \sum_{i=1}^n \bar{y}_i$$

Dlatego hipotetyczna, sumaryczna struktura finansowa w obu systemach

finansowych jest wyrażona jako $\frac{\sum_{i=1}^n \bar{y}_i}{\sum_{i=1}^n \bar{e}_i}$. Jednakże, przy tej konstrukcji,

stosunek długu do udziałów wszystkich indywidualnych banków jest ϕ razy większy w drugim systemie finansowym. Jedyna restrykcja wobec stałej ϕ wynika z warunku, że i -ty wiersz macierzy Π' sumuje się do $1 - z_i / \Phi$. Zatem ϕ nie powinno być na tyle małe, aby dla pewnych i zachodziło $1 - z_i / \phi < 0$.

Nakłada to dolną granicę na ϕ , ale nie ma granicy górnej. Można zatem

skonstruować system finansowy, gdzie sumaryczna hipotetyczna struktura finansowa jest niezmienną, ale hipotetyczna struktura finansowa indywidualnych banków może być tak wysoka jak chcemy. Pozwala to bankom na pożyczanie sobie dużych kwot, tak że ich struktura finansowa jest wyższa, ale bez zachwiania sumarycznej relacji pomiędzy sektorem bankowym i najważniejszymi wierzycielami.

Struktura finansowa całego systemu bankowego jest połączona ze strukturą indywidualnych banków w poniższy sposób, przy czym oznaczając strukturę finansową całego sektora bankowego jako L , można zapisać:

$$L = \frac{\sum_{i=1}^n y_i}{\sum_{i=1}^n e_i} = 1 + \frac{\sum_{i=1}^n e_i z_i (\lambda_1 - 1)}{\sum_{i=1}^n e_i} \quad (16)$$

gdzie wzór (16) wynika ze wzoru (13).

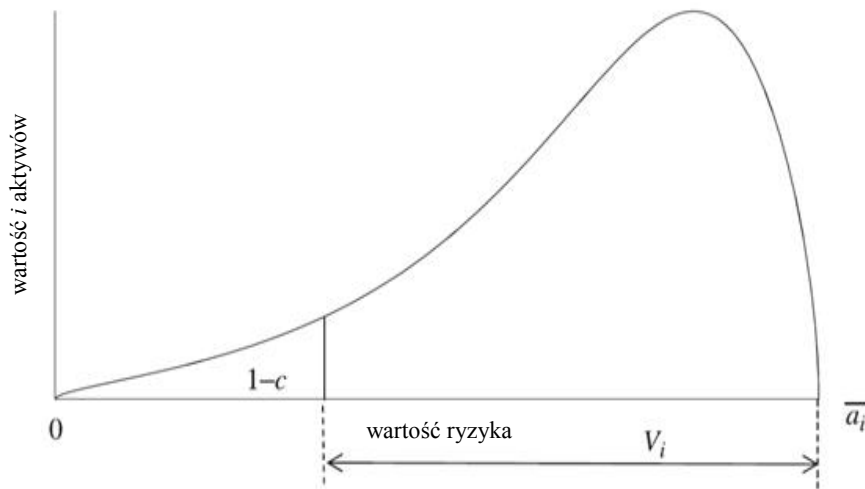
3. Wartość ryzyka

W tej części uwaga zostanie skupiona na regułach decyzyjnych stosowanych przez banki, co pozwoli na wyjaśnienie w jaki sposób udzielanie pożyczek odbiorcom końcowym zależy od poniższych parametrów motywujących ryzyko kredytowe⁹ (rys. 1). Dla banku i jego wartość ryzyka w stopniu wiarygodnym c odnosi się do wartości jego aktywów $\bar{a}_i$ i jest najmniejszą dodatnią liczbą V_i , taką że:

$$\Pr(\hat{a}_i < \bar{a}_i - V_i \leq 1 - c) \quad (17)$$

Wartość ryzyka V_i jest w przybliżeniu najgorszą stratą finansową, którą może ponieść bank, z prawdopodobieństwem mniejszym niż kryterium $1 - c$.

⁹ Na podstawie Value at Risk.



Rys. 1. Wartość ryzyka w zależności od realnych aktywów banku i

Źródło: Na podstawie H.S. Shin, op. cit., s. 320.

Dla uproszczenia zakładamy, że banki postępują zgodnie z zaleceniami wynikającymi z teorii wartości ryzyka i badają konsekwencje takiego zachowania. W szczególności bank i dąży do zrównania udziałów rynkowych e_i do wartości ryzyka V_i , zatem

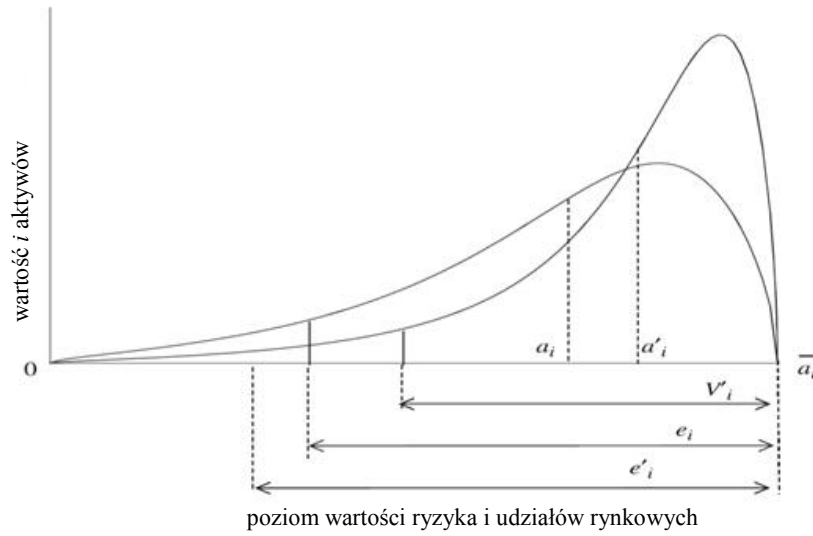
$$e_i = V_i \quad (18)$$

3.1. Zmniejszenie prawdopodobieństwa niewykonania umowy

Dla uproszczenia zakłada się, że $p_i = p$ dla wszystkich i oraz że p jest malejące. Uwzględniając funkcję dystrybuanty dla realnych spłat $\hat{y}_i$ w spłatach pożyczek banku i , otrzymujemy:

$$\Phi = \left[\frac{\Phi^{-1}(p) + \sqrt{1-\rho} \Phi^{-1}\left(\frac{z/\bar{y}_i}{\sqrt{\rho}}\right)}{\sqrt{\rho}} \right] \quad (19)$$

W przypadku kiedy wartość roszczeń międzybankowych rośnie w sumarycznym czynniku Y , spadek p pociąga za sobą zmianę stochastyczną pierwszego stopnia w gęstości realnej wartości aktywów międzybankowych $\sum_{j=1}^n \pi_{ji} \hat{x}_j$ będących w posiadaniu banku i . Wynika z tego, że istnieje stochastyczna zmiana pierwszego stopnia w gęstości realnej wartości aktywów $\hat{a}_i$ banku i (rys. 2).



Rys. 2. Wartość ryzyka i struktury finansowej

Źródło: Na podstawie: H.S. Shin, op. cit.

Rynkowa wartość aktywów przy malejącym p jest dana za pomocą a'_i , a udziały rynkowe – za pomocą e'_i . Mamy zatem $e'_i > e_i$, ponieważ wartość udziałów w terminie ostatecznym jest rosnąca $\hat{y}_i$ i istnieje zmiana stochastyczna pierwszego stopnia w $\hat{y}_i$. W tym samym czasie następuje zmniejszenie wartości ryzyka w banku i . Dzieje się tak, ponieważ wartość funkcji dystrybucyjnej obniża się wraz ze spadkiem p . Z tego powodu $(1 - c)$ wartości realnej aktywów rośnie. Wartość ryzyka jest mniejsza niż wcześniej i jest dana za pomocą V'_i , dlatego zgodnie z malejącą wartością p istnieje zależność:

$$e'_i > e_i > V'_i \quad (20)$$

To tłumaczy, dlaczego bank i ma nadwyżkę udziałów w takim sensie, że ilość jego udziałów rynkowych jest wyższa niż ilość udziałów wymagana przez wartość ryzyka. Nadwyżka udziałów może być ograniczona przez płacenie dywidendy pozostałym udziałowcom, lub przez ponowne wykupienie udziałów poprzez wyemitowanie większego zadłużenia. W praktyce banki redukują jednak ilość udziałów poprzez zwiększenie wielkości bilansów finansowych, niż przez spłacanie nadwyżek¹⁰. Zgodnie z powyższym zakłada się, że jeśli bank i ma nadwyżkę udziałów, rozszerzy swój bilans poprzez zwiększenie hipotetycznej wartości długu $\bar{x}_i$ i użyje zysków, aby nabyć więcej aktywów. Można zatem przyjąć, że kiedy $e'_i > V'_i$ po zmniejszeniu p , bank i zwiększa wartość nominalną swojego zadłużenia x_i .

Kiedy banki zwiększają zadłużenie, nabywają aktywa z zyskiem. Macierz roszczeń międzybankowych Π ulegnie wtedy zmianie. Zakładając, że nowa macierz roszczeń międzybankowych jest dana jako Π^* , wtedy odpowiednio wartości nominalne i profil wartości rynkowej długu oznacza się jako x^* . Kiedy rośnie wartość nominalna długu banku i , wartość rynkowa długu rośnie dla wszystkich banków¹¹:

$$x_i^* \geq x_i \quad (21)$$

Zakładamy również, że jeśli banki zwiększają swoje pożyczki ze względu na pojawienie się nadwyżki udziałów, będą szukać nowych źródeł finansowania. Jeśli innowacje finansowe w postaci sekurytyzacji są dostępne, banki mogą wykorzystać je przez pożyczanie od sektora zewnętrznych wierzycieli. Pozwala to na stworzenie kolejnego założenia – kiedy banki zwiększają swój hipotetyczny dług w odpowiedzi na spadek p , proporcja funduszy pozyskanych z sektora zewnętrznego jest niemalejąca. To założenie nakłada ograniczenie na nową macierz roszczeń międzybankowych Π^* , tak że suma jej i -tego rzędu nie może być większa niż suma odpowiadającego rzędu i w macierzy Π :

$$(I - \Pi^*)u \geq (I - \Pi)u \quad (22)$$

Kiedy spada p , wartość sumaryczna pożyczek do odbiorców końcowych rośnie, zarówno w wartościach hipotetycznych, jak i rynkowych. Tożsamości arkusza finansowego w wartości nominalnej są następujące:

$$\begin{aligned} \bar{y} &= \bar{e} + \bar{x}(I - \Pi) \\ \bar{y}^* &= \bar{e}^* + \bar{x}^*(I - \Pi^*) \end{aligned} \quad (23)$$

¹⁰ T. Adrian, H.S. Shin, *Liquidity and Leverage*, „Journal of Financial Intermediation” 2007, <http://www.princeton.edu/~hsshin/working.htm>; Eidem, *Financial Intermediary Leverage and Value at Risk*, Federal Reserve Bank of New York and Princeton University 2008, <http://www.princeton.edu/~hsshin/working.htm>.

¹¹ L. Eisenberg, T.H. Noe, *Systemic Risk in Financial Systems*, „Management Science” 2001, Vol. 47.

gdzie * wskazuje zmienne po zmianie. Wartość nominalna udziałów pozostaje niezmieniona ($\bar{e}^* = \bar{e}$), tak że zmiana sumarycznej hipotetycznej wartości pożyczek jest dana:

$$(\bar{y}^* - \bar{y})u = (\bar{x}^* - \bar{x})(I - \Pi)u + \bar{x}^*(\Pi - \Pi^*)u \quad (24)$$

Pierwszy wyraz po prawej stronie jest dodatni z założenia, że bank reaguje na nadwyżkę udziałów poprzez rozszerzenie swojego bilansu, podczas gdy drugi wyraz po prawej stronie jest dodatni z równania (22) o rosnącej proporcji funduszy pochodzących z sektora zewnętrznego. Dlatego $(\bar{y}^* - \bar{y})u > 0$, a zatem hipotetyczna całkowita wartość pożyczek do odbiorców końcowych rośnie.

Argumentem za wzrostem wartości rynkowej pożyczek jest to, że spadek p jest taki sam. Tożsamość wartości rynkowych arkuszy finansowych przed i po zmianie jest następująca:

$$\begin{aligned} y &= e + x(I - \Pi) \\ y^* &= e^* + x^*(I - \Pi^*) \end{aligned} \quad (25)$$

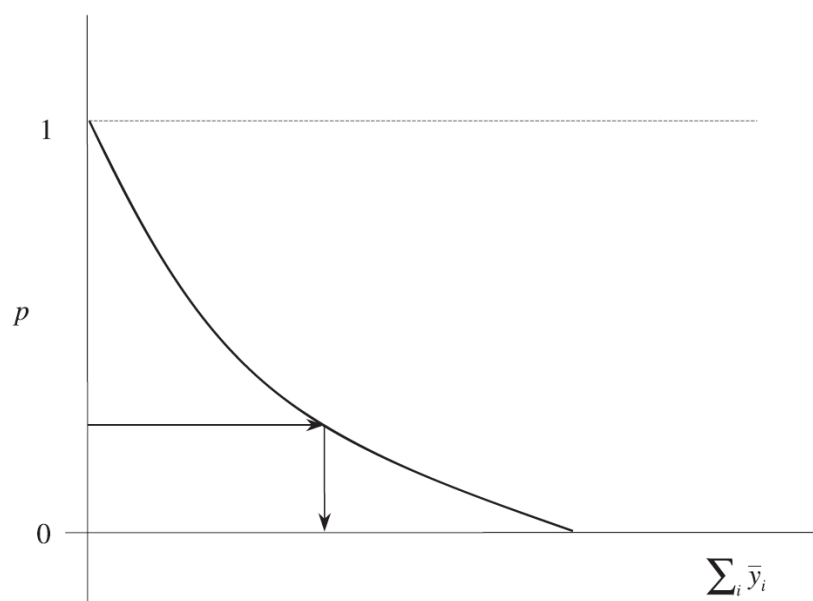
Zmiana wartości rynkowej pożyczek do odbiorców końcowych jest wyrażona:

$$(y^* - y)u = (e^* - e)u + (x^* - x)(I - \Pi)u + x^*(\Pi - \Pi^*)u \quad (26)$$

Równanie (26) różni się od analogicznego dla wartości nominalnych, tym że bilanse banków uwzględniają teraz główny zysk na swoich pożyczkowych portfelach papierów wartościowych poprzez $(e^* - e)u$, gdzie $e^* = e'$ i gdzie e' jest wartością podaną w (20). Rosnąca ilość udziałów jest dodatkowym źródłem finansowania, kiedy pożyczki są przewartościowywane na wartości rynkowe. Wszystkie trzy wyrazy po prawej stronie (26) są dodatnie, zatem $(y^* - y)u > 0$.

4. Boom pożyczkowy

Na podstawie rezultatów wyprowadzonych do tej pory można naszkicować scenariusz boomu pożyczkowego. Pierwszym składnikiem jest związek pomiędzy prawdopodobieństwem zaniechania p i sumaryczną ilością pożyczek dla odbiorców końcowych. Rysunek 3 pokazuje negatywny związek między całkowitą (hipotetyczną) ilością pożyczek i wartością p , gdzie całkowita ilość pożyczek jest opisana na osi poziomej. Strzałki pokazują, że dla każdego poziomu p , istnieje powiązany poziom całkowitej ilości udzielanych pożyczek $\sum_i \bar{y}_i$.



Rys. 3. Sumaryczna ilość pożyczek w stosunku do współczynnika zaniedbań p

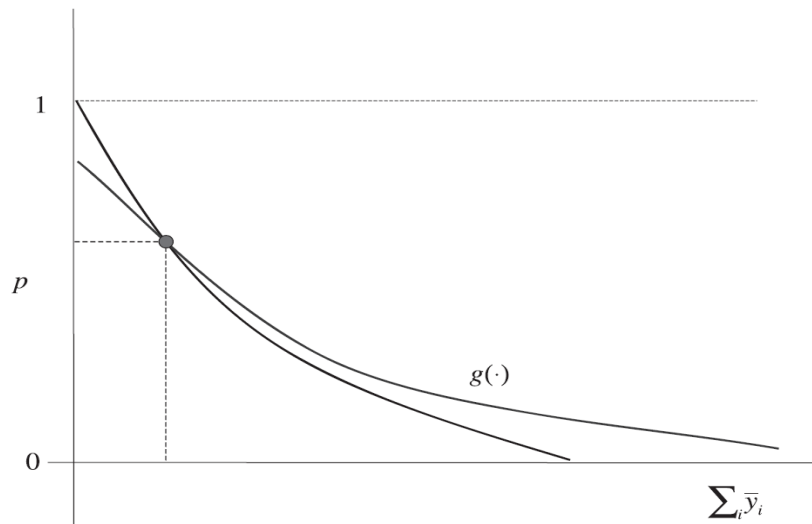
Źródło: Na podstawie: H.S. Shin, op. cit., s. 324.

Zakładając dalej, że istnieje makroekonomiczne sprzężenie, począwszy od całkowitej ilości udzielanych pożyczek do prawdopodobieństwa zaniedbań, można wtedy oczekiwać przyrostów, które skutkują wzajemnym oddziaływaniem pomiędzy wzmacnianymi bilansami finansowymi a rosnącą ilością pożyczek¹². W celu zilustrowania rosnącego zapotrzebowania na pożyczki, przyjmuje się malejące odwzorowanie g , które opisuje sumaryczną ilość pożyczek $\sum_i \bar{y}_i$ do prawdopodobieństwa zaniedbań p .

Uwzględniając funkcję $q(\cdot)$ (rysunek 4 nakłada funkcję $q(\cdot)$ na rysunek 3), można rozważyć scenariusz, w którym sekurytyzacja zmieni proporcję pomiędzy wewnętrznym i zewnętrznym finansowaniem na korzyść zwiększenia ilości funduszy pochodzących z zewnątrz. Można zinterpretować ten scenariusz jako spadek wpisów do międzybankowej macierzy roszczeń Π , tak że fundusze pozyskiwane z sektora pozabankowego proporcjonalnie rosną.

¹² T. Adrian, H.S. Shin, *Financial Intermediaries, Financial Stability and Monetary Policy*, The Federal Reserve Bank of Kansas City Symposium at Jackson Hole 2008, <http://www.princeton.edu/~hsshin/working.htm>.

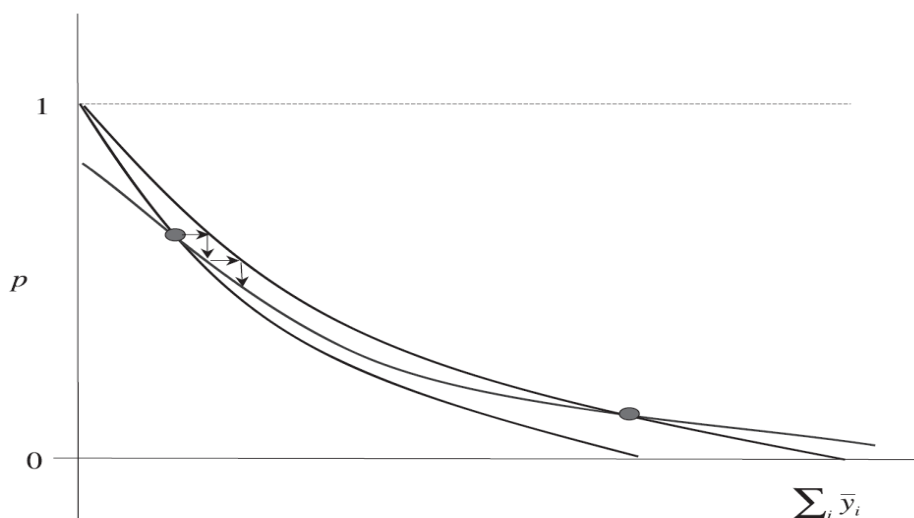
Jak już zostało omówione w poprzednich rozdziałach, rozwój gospodarczy zwiększa sumaryczną ilość pożyczek do odbiorców końcowych, nawet jeśli struktura finansowa indywidualnych banków (i ich wartość ryzyka) pozostają niezmienione. Na wykresach widać, że zmiana w kierunku większego wykorzystania funduszy zewnętrznych, może być reprezentowana jako zmiana prawej strony równania podaży kredytowej. Rysunek 5 pokazuje zmiany w podaży kredytowej oraz jej rezultaty.



Rys. 4. Punkt początkowy/inicjujący

Źródło: Ibid.

Początkowy wstrząs wynikający z większego wykorzystania funduszy zewnętrznych skutkuje prawostronną zmianą w wykresie podaży kredytów. Nowy punkt przecięcia się jest na dole po prawej stronie w stosunku do punkt początkowego, połączony z mniejszym prawdopodobieństwem zaniedbań p i większą sumaryczną ilością pożyczek.



Rys. 5. Boom pożyczkowy

Źródło: Ibid.

Zmiana stanu początkowego jest zmianą prawostronną w krzywej podaży kredytowej, która skutkuje większą sumaryczną ilością pożyczek dla stałego p , jednakże reakcją w skali makro na większą podaż kredytów, jest obniżenie ryzyka zaniechań p . To nastawienie jest pokazane przez pierwszą skierowaną w dół strzałkę na rysunku 5. Spadek p skutkuje zwiększeniem ilości przyznawanych pożyczek. Zwiększenie ilości pożyczek znów obniża p itd. Nowy punkt przecięcia na rysunku 5 jest połączony zarówno z wyraźnie niższym prawdopodobieństwem zaniechań, jak i z większą ilością pożyczek.

Podsumowanie

Znaczenie sekurytyzacji dla stabilności finansowej wywodzi się ze zdolności systemu bankowego do zwiększenia podaży kredytów dla odbiorców końcowych. Kiedy istnieje spadek w ryzykowności aktywów podstawowych (w tym przypadku poprzez spadek parametru p), wzrasta możliwość podejmowania ryzyka przez system bankowy.

Idea głosząca, że zmiany w bilansach finansowych pożyczkodawców są ważne w determinowaniu podaży kredytów, pojawiła się już we wcześniejszej literaturze, podkreślając płynną strukturę arkuszy bankowych lub efekt poduszki amortyzującej w kapitale regulacyjnym banku.

Zgodnie z zaprezentowanym scenariuszem w tym artykule, kryzys hipoteczny ma swój początek w rosnącej podaży kredytowej lub jednoznacznie w potrzebie znalezienia aktywów, służących do wypełnienia poszerzonych bilansów finansowych, co przekłada się na obniżanie bankowych standardów pożyczkowych. W ten sposób możliwe staje się wytłumaczenie dwóch cech tego kryzysu. Dlaczego duże, zaawansowane systemy pośrednictwa finansowego wciąż pożyczają pieniądze mało wiarygodnym pożyczkobiorcom, i po drugie, dlaczego zaawansowane systemy bankowe trzymały złe pożyczki w swoich sprawozdaniach finansowych, a nie przekazywały ich dalej niczego niepodważającym inwestorom. Obydwa fakty można wytłumaczyć potrzebą „łatania dziur” we własnych, wciąż rozszerzanych zestawieniach finansowych.

ESTIMATION OF SECURITIZATION AS A PROCESS INFLUENCING FINANCIAL STABILITY OF THE BANKING SYSTEM

Summary

The article presents considerations concerning securitization in relation to the concept of financial stability of the banking system. The very process of securitization has been widely used in developed countries, both in banking and insurance. After the recent financial crisis this process has slowed down and is even indicated as a source of collapse. Based on theoretical considerations, literature studies, and specific operational scenarios, has been made an attempt to verify this hypothesis. For this purpose, special attention is paid to increasing the supply of credit and appropriate accounting policies with regard to credit risk. The next section highlights the level of debt, market value and risk in the financial system of the bank. At the end, lending boom is characterized, and are given the conclusions coming out of studied accounting scenarios.