

**PROGNOSTYCZNE UWARUNKOWANIA
RYZYKA
GOSPODARCZEGO I SPOŁECZNEGO**

Studia Ekonomiczne

ZESZYTY NAUKOWE

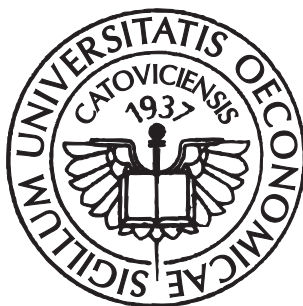
WYDZIAŁOWE

UNIWERSYTETU EKONOMICZNEGO

W KATOWICACH

PROGNOSTYCZNE UWARUNKOWANIA RYZYKA GOSPODARCZEGO I SPOŁECZNEGO

**Redaktor naukowy
Włodzimierz Szkutnik**



Katowice 2013

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, Gwo-Hsiung Tzeng,
Zdeněk Mikoláš, Marian Noga, Bronisław Micherda, Miloš Král

Recenzenci

Jerzy Wiśniewski
Tadeusz Stanisław

Redaktor

Barbara Cebo

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego
w Katowicach 2013

ISBN 978-83-7875-060-4

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.: +48 32 257-76-35, faks: +48 32 257-76-43
www.wydawnictwo.ue.katowice.pl e-mail: wydawnictwo@ue.katowice.pl

SPIS TREŚCI

WSTĘP	9
Włodzimierz Szkutnik	
CHAOTYCZNE REAKCJE RYNKÓW FINANSOWYCH – ASPEKT PROBABILISTYCZNY WYCENY I ZABEZPIECZEŃ PŁATNICZYCH NA RYNKU KAPITAŁOWYM	11
Summary	28
Konstancja Poradowska	
MODELE SUBIEKTYWNE W KONSTRUKCJI PROGNOZ DŁUGOOKRESOWYCH	29
Summary	43
Włodzimierz Szkutnik	
STATYSTYCZNA NIEOKREŚLONOŚĆ W WYCENIE CHARAKTERYSTYK RYNKÓW FINANSOWYCH	45
Summary	62
Jerzy Zemke	
RYZYSKO W ASPEKCIE ZARZĄDZANIA W ZRÓŻNICOWANYM OTOCZENIU SPOŁECZNO-GOSPODARCZYM	63
Summary	76
Maria Balcerowicz-Szkutnik	
UWARUNKOWANIA POZIOMU BEZROBOCIA WYBRANYCH PAŃSTW UE – – ANALIZA STATYSTYCZNA	77
Summary	86
Anna Sączewska-Piotrowska	
PROGNOSTYCZNY WARIANT UBÓSTWA DLA GOSPODARSTW DOMOWYCH MAKROREGIONU POŁUDNIOWEGO	87
Summary	98

Mirosław Wójciak	
METODY BUDOWY DŁUGOTERMINOWYCH PROGNOZ PRZEDZIAŁOWYCH ROZWOJU NOWYCH ZJAWISK	99
Summary	113
Alicja Wolny-Dominiak	
MODELOWANIE LICZBY SZKÓD W UBEZPIECZENIACH KOMUNIKACYJNYCH W PRZYPADKU WYSTĘPOWANIA DUŻEJ LICZBY ZER, Z WYKORZYSTANIEM PROCEDURY KROSWALIDACJI.....	115
Summary	129
Tadeusz Czernik	
WPŁYW NIEPEWNOŚCI OSZACOWANIA ZMIENNOŚCI NA CENĘ INSTRUMENTÓW POCHODNYCH.....	131
Summary	142
Monika Dyduch	
BANKOWE PAPIERY WARTOŚCIOWE STRUKTURYZOWANE.....	143
Summary	164
Iwona Dittmann	
PODOBIEŃSTWO ZMIAN ŚREDNICH CEN TRANSAKCYJNYCH 1 m ² POWIERZCHNI MIESZKAŃ W WYBRANYCH MIASTACH WOJEWÓDZTWA ŚLĄSKIEGO.....	165
Summary	182
Daniel Iskra	
WARTOŚĆ ZAGROŻONA OPCJI EUROPEJSKICH SZACOWANA PRZEDZIAŁOWO. SYMULACJE.....	183
Summary	192
Jan Acedański	
KRYTERIA WYBORU DYNAMICZNYCH MODELI CZYNNIKOWYCH DLA CELÓW PROGNOSTYCZNYCH	193
Summary	216
Maciej Pichura	
ANALIZA WPŁYWU KOSZTÓW TRANSAKCYJNYCH NA OCENĘ EFEKTYWNOŚCI WYBRANEJ STRATEGII INWESTYCYJNEJ.....	217
Summary	231

Agnieszka Przybylska-Mazur

REGUŁY POLITYKI PIENIĘŻNEJ A PROGNOZOWANIE

WSKAŹNIKA INFLACJI 233

Summary 244

WSTĘP

W Studiach Ekonomicznych pt. *Prognostyczne uwarunkowania ryzyka gospodarczego i społecznego* autorzy podjęli tematykę uwarunkowań ryzyka w szerokim spektrum zastosowań gospodarczych i społecznych. Uwaga autorów zwrócona została na aspekt prognostyczny prowadzonych analiz.

W obszarze analiz ryzyka rynku kapitałowego ich tematyka objęła reakcje rynków finansowych na nielosowe zaburzenia obserwowane w pewnych uwarunkowaniach o charakterze chaotycznych procedur (W. Szkutnik). Uwzględniony został charakter nieokreśloności statystycznej obserwowany w wycenie charakterystyk rynków finansowych, co oddaje w pewnym sensie nie-realność założeń pierwotnie skonstruowanego modelu portfelowej analizy Markowitza (W. Szkutnik). W analizach tych uwzględnia się aspekt probabilistyczny wyceny i zabezpieczeń płatniczych na rynku kapitałowym, z czym łączy się pewien wątek prognozowania procesów tam obserwowanych. Z tym zakresem rozważań łączy się badanie wpływu niepewności oszacowania zmienności na cenę instrumentów pochodnych (T. Czernik) oraz wartość zagrożona europejskich opcji szacowana przedziałowo (D. Iskra). Tematykę ryzyka kapitałowego podejmuje się też w analizie wpływu kosztów transakcyjnych na ocenę efektywności inwestycji kapitałowej (M. Pichura). Z tematami ryzyka kapitałowego sąsiaduje problematyka, jaką generują formy bankowego ryzyka, którego aspekt zabezpieczania wyrażają i są jego emanacją bankowe papiery wartościowe strukturyzowane (M. Dyduch). Te instrumenty, których konstrukcja na niwie nauki jest jeszcze mało rozwinięta zasługują na większe uznanie badaczy niż mogłoby się to wydawać. Propozycje nowych instrumentów, czy to giełdowych czy bankowych, powinny wynikać z propozycji środowiska badaczy i naukowców. Nie jest to w żadnym przypadku domena praktyków bankowości, a bynajmniej nie powinno tak być w realnej sferze działalności bankowej. Pewne paralele ze środowiskiem finansów widoczne są w opracowaniu analizującym wybrane charakterystyki z rynku nieruchomości (I. Dittmann). Rynek nieruchomości i rynek finansów unifikują się z coraz większym natężeniem, a zatem ryzyko na rynku nieruchomości staje się ekwiwalentne do ryzyka finansowego.

W innym obszarze badań prognostycznych zachowań o charakterze długofalowym umiejscowione zostały opracowania traktujące o metodach budowy prognoz przedziałowych działań, których realizacja może nastąpić w odległej przyszłości, co dotyczy np. nowych generatorów energii (M. Wójciak) oraz subiektywizmu w konstrukcji modeli prognoz długofalowych (K. Poradowska).

Zjawiska i procesy, które mogą być obciążone także ryzykiem o charakterze społecznym stały się przedmiotem badań w opracowaniach, których tematyka odnosi się do uwarunkowania poziomu bezrobocia w przekroju pewnych państw UE (M. Balcerowicz-Szkutnik), a także prognostycznego wariantu ubóstwa (A. Sączewska-Piotrowska) dla makroregionu południowego Polski. Tematyka ta styka się dosłownie z procesami ekonomicznymi i jej badanie w kategoriach ryzyka zasługuje na traktowanie jej z całą dozą znaczenia w rozstrzygnięciu niepewności bytu społeczeństw.

Ryzyko ubezpieczeniowe najbardziej wyrażające jego istotę rozpatrzono w analizie modelowej liczby szkód w ubezpieczeniach komunikacyjnych (A. Wolny-Dominiak). Zwrócono tu uwagę na nieprzystawalność w pewnych wypadkach modelowania liczby szkód przez rozkład Poissona i podjęto próbę wykorzystania procedury krosvalidacji, gdy z liczbą szkód pojawia się występowanie dużej liczby zer w ubezpieczeniach komunikacyjnych.

Istota celów prognostycznych jest rozwinięta w bardzo wyrafinowanym stylu z pozycji kryteriów wyboru modeli czynnikowych (J. Acedański). Propozycje autora naszkicowane w proponowanym ujęciu są wzbogacone szczegółowo omówionymi badaniami empirycznymi.

W ostatniej propozycji uwarunkowania ryzyka w profilu zarządzania w zróżnicowanym otoczeniu społeczno-gospodarczym, a więc stricte łączącym tytułowe jego formy, zawarte zostały w opracowaniu (J. Zemke), w wyprofilowany sposób traktującym o formach, możliwościach i znaczeniu zarządzania ryzykiem, co przekłada się na podejmowanie decyzji menedżerskich w szerokim obszarze aplikacji.

Niepełność tematyki zarysowanej szeroko w temacie niniejszego Zeszytu Naukowego jest bez wątpienia inspiracją dla dalszych badań ukierunkowanych na bardziej praktyczne pole odkrywczych inspiracji w zakresie ryzyka gospodarczego i społecznego. Tematyka ta ma duże możliwości rozwojowe i może wyrażać wiele innowacyjnych formuł teoretycznych i praktycznych.

Włodzimierz Szkutnik

Włodzimierz Szkutnik

Uniwersytet Ekonomiczny w Katowicach

CHAOTYCZNE REAKCJE RYNKÓW FINANSOWYCH – – ASPEKT PROBABILISTYCZNY WYCENY I ZABEZPIECZEŃ PŁATNICZYCH NA RYNKU KAPITAŁOWYM

Wprowadzenie

Właściwe dla prowadzonych rozważań będą modele finansowych kalkulacji na zupełnych i niezupełnych rynkach z zastosowaniem niesamofinansujących się strategii. W takich wypadkach naturalnym ujęciem zagadnienia jest odejście od typowego założenia dotyczącego stochastycznej struktury cen akcji przy modelowaniu stóp zwrotu (ich logarytmów) i rozpatrzenie wariantowego przypadku, w którym nie czyni się założeń dotyczących rozkładu normalnego logarytmów tych stóp. Wymaga to jednak egzemplifikacji modelu finansowego rynku w warunkach statystycznej nieokreśloności, co odpowiada zadaniu modelowania racjonalnego zachowania inwestora. W wypadku niezupełnych rynków rozważone będą kalkulacje opcyjnie i zadania minimalizacji ryzyka.

1. Zupełny rynek i strategie arbitrażowe

W analizach portfelowych uwzględniających ryzyko inwestycyjne rozważa się w formalnym ujęciu model finansowego rynku i inwestycyjne strategie umożliwiające opis ewolucji papierów wartościowych na finansowym rynku. Wystarcza wtedy przyjąć, że jego postać wyrażona jest przez dwa dyskretne stochastyczne równania opisujące aktywa pozbawione ryzyka i aktywa ryzykowne.

Taki model można zapisać w postaci:

$$\Delta B_n = r_n \cdot B_{n-1} \quad , \quad \Delta S_n = \rho_n \cdot S_{n-1}$$

Dyskretyzacja stochastyczna tych równań wymaga, aby przestrzeń probabilistyczna będąca opisem losowego rozkładu wartości aktywów była generowana przez skończony zbiór zdarzeń elementarnych Ω , przeliczalną rodzinę zdarzeń losowych odpowiadających borelowskiej algebrze zdarzeń $F = F_N, N \in C_+$ – zbiór całkowitych liczb dodatnich. Wprowadzenie oznaczenia dla sum pierwszych n wyrazów (stochastycznych) ciągów $(\rho_n)_{n \in Z_+}$ i $(r_n)_{n \in Z_+}$

$$T_n = \sum_{k=0}^n r_k, \quad W_n = \sum_{k=0}^n \rho_k$$

pozwala sprowadzić równania modelu rynku dyskretnego do postaci stochastycznie eksponentcjanej:

$$B_n = B_0 \cdot \varepsilon_n(T), \quad S_n = S_0 \cdot \varepsilon_n(W) \quad (1)$$

Przy takim opisie rynku zakłada się, że $F_n = \sigma\{S_0, \dots, S_n\}$, tj. F_n jest minimalną σ -algebrą, względem której mierzalne są zdarzenia $S_0, \dots, S_n$.

Rynek określony w równaniach (1) nazywany jest powszechnie (B, S) -rynkiem. W dalszej części dla zapobieżenia niepotrzebnym technicznym trudnościom wywody będą prowadzone dla przypadku jednowymiarowego S_n , i w tym wypadku, jak wynika ze znanych z literatury wyników, najbardziej treściwym modelem (zupełnym) (B, S) -rynku jest dwumianowy model, jednak wiele wyników jest także adekwatnych w wielowymiarowym wariancie. Okazuje się także, że można uzyskać ogólny model (B, S) -rynku zakładając tylko dodatnią wartość cen aktywów B i S . W tym wypadku pojawia się multiplikatywna forma dla B i S , która znowu prowadzi do rozpatrywanego tutaj modelu (1).

W analizie inwestycji kapitałowych stosowane jest pojęcie tzw. inwestycyjnej strategii, przez którą rozumiany jest stochastyczny dwuwymiarowy ciąg $\pi = (\pi_n(\beta_n, \gamma_n))_{n \in Z_+}$. Elementy $\beta_n \in F_n, \gamma_n \in F_{n-1}$ tego ciągu interpretowane są jako liczby aktywów bez ryzykownych B i z ryzykiem S w momencie czasu $n \in Z_+$. Ponadto istotne jest także pojęcie „kapitału” portfela π , którym w myśl przyjmowanego określenia jest stochastyczny ciąg $X^\pi = (X_n^\pi)_{n \in Z_+}$, gdzie $X_n^\pi = \beta_n \cdot B_n + \gamma_n \cdot S_n$.

Na podstawie tych pojęć można wprowadzić klasę samofinansujących się portfeli, przez którą będziemy rozumieć klasę portfeli π oznaczaną przez SF , spełniającą warunek:

$$B_{n-1} \cdot \Delta \beta_n + S_{n-1} \cdot \Delta \gamma_n = 0 \quad (2)$$

W tym kontekście można zauważyć, że kapitał samofinansującego się portfela π wyrażony w postaci:

$$X_n^\pi = X_0^\pi + \sum_{k=0}^n (\beta_k \cdot \Delta B_k + \gamma_k \cdot \Delta S_k)$$

jest równoważny warunkowi samofinansującego się portfela (2), gdzie $\Delta B_0 = \Delta S_0 = 0$.

Istotna w stosowaniu strategii inwestycyjnych jest dopuszczalność arbitrażu na rynku akcji. Dlatego w klasie portfeli samofinansujących się SF można wyróżnić te portfele π , które realizują arbitrażową możliwość na rynku akcji w następującym znaczeniu:

$$X_0^\pi = 0, X_n^\pi \geq 0$$

dla $n \leq N$ (z prawdopodobieństwem P – prawie na pewno (P – p.n.p) i $X_N^\pi > 0$ z dodatnim prawdopodobieństwem).

Ekonomiczna treść, która tu się przejawia wynika z określenia występowania arbitrażu na rynku. Polega to na możliwości pojawienia zysku z inwestycji bez ryzyka, gdy na rynku występuje arbitrażowy portfel. Klasę takich portfeli oznaczmy przez SF_{arb} i rynek nazwiemy arbitrażowym lub bez arbitrażu, gdy odpowiednio w klasie SF_{arb} występuje chociaż jeden arbitrażowy portfel lub gdy takiego portfela nie ma.

W dalszej części opracowania będą rozważane finansowe kalkulacje na pełnym rynku, przy niesamofinansujących się strategiach. Prowadzenie takich wywodów wymaga jednak pewnych ścisłych określeń z zakresu probabilistyki i wyprowadzonych na ich bazie własności. Wydaje się, że dla pełnego zrozumienia dalszych rozważań niezbędne jest omówienie chociaż podstawowych pojęć z tego zakresu i podanie znanych wyników z teorii procesów stochastycznych.

2. Martyngałowe miary i arbitraż

W probabilistycie miarę prawdopodobieństwa P^* , równoważną P , nazywa się miarą martyngałową lub neutralną względem ryzyka, jeśli względem miary P^* stochastyczny ciąg:

$$(S_n/B_n)_{n \leq N}$$

jest martyngałem. Oznacza to stałość wartości oczekiwanych względem danej miary probabilistycznej [4] dla wyrazów powyższego ciągu. Miar takich może być cała klasa P^* .

Kryterium martyngalności miary P'

Jeśli w modelu rynku (1) stochastyczny ciąg $(r_n)_{n \leq N}$ jest prognozowalny i $r_n \rightarrow -1$, wtedy względem miary P' jednocześnie są martyngalami:

$$R_n = \frac{S_n}{R_n} \quad \text{oraz} \quad (\sum_{k=0}^n (\rho_k - r_k))_{n \leq N}$$

Kryterium to wynika głównie z własności stochastycznej wykładniczości.

Martyngalowa miara P^* względem miar równoważnych P

W tej sytuacji zupełnie naturalnie pojawia się problem poszukiwania martyngalowej miary P^* wśród miar równoważnych P . Oznaczając w tym celu odpowiadającą lokalną gęstość przez:

$$(Z_n)_{n \leq N}$$

z kryterium martyngalności miary względem P^* wynika, że:

$$R \text{ jest martyngalem} \Leftrightarrow (V - U)$$

Dla przykładu i upraszczając kontekst powyższego formalnego ujęcia przyjmujemy, że V jest martyngalem już względem wyjściowej miary P . Wtedy z twierdzenia Girsanowa [4] wynika, że:

$$V_n^* = V_n - \sum_{k \leq n} E(Z_{k-1}^{-1} \cdot Z_k \cdot \Delta V_k / F_{k-1})$$

jest martyngalem względem miary P^* . Konsekwencją tego jest sposób wyboru miary P^* , która powinna być wybrana w taki sposób, aby odpowiadająca jej gęstość czyniła zadość relacji:

$$\Delta U_n = E(Z_{n-1}^{-1} \cdot Z_n \cdot \Delta V_n / F_{n-1})$$

Przypadek ten ma naturalne uogólnienie.

Wprowadzone powyżej pojęcia i podane wnioski prowadzą do stwierdzenia ścisłego związku arbitrażowego rynku, będącego ekonomiczną kategorią w aspekcie finansowego postrzegania inwestycji na rynku kapitałowym oraz martyngalowej miary. Zachodzi bowiem równoważność między istnieniem miary martyngalowej P^* wśród miar równoważnych P a istnieniem arbitrażowego portfela SF_{arb} , gdy w modelu rynku (1) o ciągu stóp zwrotu

r_n ($r_n > -1, n \leq N$) założy się deterministyczną naturę [8]. W nietrywialnym dowodzie tego faktu występują dwa zbiory zmiennych losowych $\xi = \xi(\omega)$ określonych na (Ω, F) , które będą jeszcze wykorzystane w dalszej części rozważań:

$$\sum_0 = \{ \xi \in R : \exists \pi \in SF, X_0^\pi = 0 \wedge X_N^\pi = \xi \}$$

$$\sum_1 = \{ \xi \geq 0 : E\xi \geq 1 \}$$

Okazuje się, że przy nieistnieniu arbitrażowego portfela wynika, że zbiory te nie mają wspólnych elementów.

3. Urealnienie kierunków modelu rynku

Badanie rynku (1) może być prowadzone w kierunku bardziej realnej sytuacji, kiedy zmiany portfela są stowarzyszone albo z przyływem, albo z odpływem kapitału. Modelowanie takiej sytuacji na rynku zostanie przeprowadzone w obecności pewnego stochastycznego ciągu $G = (G_n)$ oraz całej klasy takich strategii $\pi = (\beta_n, \gamma_n)$, które będziemy nazywać G -finansującymi, a ich klasę będziemy oznaczać przez GF . Właściwość charakteryzująca tę klasę uwzględnia odpływy i przyływy kapitału z i do portfela, co wyraża równanie:

$$B_{n-1} \cdot \Delta\beta_n + S_{n-1} \cdot \Delta\gamma_n = -\Delta G_n \quad (3)$$

gdzie:

$$G_n = \sum_{k=1}^n \Delta G_k, \quad G_0 = 0$$

Należy przyjąć, że jeśli $\Delta G \geq 0$ (odpowiednio $\Delta G \leq 0$), to G -finansującą strategię π nazywa się strategią zapotrzebowania (odpowiednio strategią z refinansowaniem lub inwestowaniem). Z powyższego wynika, że na podstawie (2) samofinansowanie oznacza 0-finansowalność.

Odpowiednio do równania (3), które nazywa się równaniem bilansowym [8] dla kapitału X^* strategii $\pi \in GF$ mamy zależności:

$$\begin{aligned} X_n^\pi &= \beta_n \cdot B_n + \gamma_n \cdot S_n \\ X_{n-1}^\pi &= \beta_n \cdot B_{n-1} + \gamma_n \cdot S_{n-1} - \Delta G_n \end{aligned} \quad (4)$$

Stąd dla dowolnej strategii π z klasy samofinansujących się portfeli SF z równania (4) otrzymuje się:

$$\Delta X_n^\pi = r_n \cdot X_{n-1}^\pi + \gamma_n \cdot S_{n-1} \cdot (\rho_n - r_n) - (1 + r_n) \cdot \Delta G_n$$

To stochastyczne niejednorodne i liniowe równanie ma rozwiązanie:

$$X_n^\pi = Exp_n(U) \cdot \{X_0^\pi + \sum_{k=1}^n \mathcal{E}_k^{-1}(U) \cdot \gamma_k \cdot S_{k-1}(\rho_k - r_k) - \sum_{k=1}^n \mathcal{E}_{k-1}^{-1}(U) \cdot \Delta G_k\} \quad (5)$$

Oznaczając:

$$M_n^\pi = X_0^\pi + \sum_{k=1}^n \mathcal{E}(U) \cdot \gamma_k \cdot S_{k-1}(\rho_k - r_k)$$

$$G_n^{exp} = \sum_{k=1}^n \mathcal{E}_{k-1}^{-1}(U) \cdot \Delta G_k, \quad G_0^\mathcal{E} = 0$$

ze wzoru (5) otrzymamy nową postać tego równania:

$$\mathcal{E}_k^{-1}(U) \cdot X_n^\pi - G_n^\mathcal{E} \quad (6)$$

Ze stwierdzonej już wcześniej własności o równoważności między istnieniem miary martyngałowej $\mathbf{P}^*$ wśród miar równoważnych $\mathbf{P}$ a istnieniem arbitrażowego portfela SF_{arb} , gdy w modelu rynku (1) o ciągu stóp zwrotu r_n ($r_n > -1, n \leq N$) założy się deterministyczny charakter, wynika w szczególności, że względem martyngałowej miary $\mathbf{P}^*$ wielkość:

$$\mathcal{E}_n(U) \cdot X_n^\pi$$

jest martyngałem, jeśli G jest martyngałem.

Jako wniosek z (6) otrzymuje się, że:

$$E^* \mathcal{E}_k^{-1}(U) \cdot X_n^\pi = X_0^\pi - \sum_{k=1}^n E^* \mathcal{E}_{k-1}^{-1}(U) \cdot \Delta G_k$$

Dla (B, S) -ryнку (1) z zadany m płatniczym rygorem (f, N) względem europejskiej lub amerykańskiej opcji zachowane zostają określenia wynikające z określenia płatniczych reguł i inwestycyjnego kosztu dla opcji europejskich i odpowiednie własności dla opcji amerykańskich. Poniżej krótko wyjaśnimy założenia specyfikujące dodatkowe właściwości samofinansującego się portfela spełniającego wprowadzone właściwości.

Zobowiązania płatnicze i opcje europejskiego rodzaju

Rozpatrując finansowy rynek (B, S) uwzględnia się tzw. zobowiązania płatnicze z datą wygaśnięcia N , przez które rozumie się parę (f, N) , gdzie f jest F_N mierzalną nieujemną losową wielkością. Uczestnik rynku, który po-

winien wygasić płatnicze zobowiązanie zmuszony jest tak zorganizować swoją inwestycyjną działalność, aby odpowiedni portfel inwestycyjny π dostarczył kapitał $X_N^\pi \geq f$.

Procedura skonstruowania takiego portfela, prowadząca do zrealizowania zobowiązania płatniczego, nazywana jest hedgingiem tego zobowiązania, a sam portfel – hedgingowanym portfelem. Charakter płatniczych zobowiązań może być dostatecznie dowolny w ramach dopuszczalnych procedur dozwolonych na rynku akcji. W tym aspekcie jedno z ważniejszych zadań wynikających z hedgingowania płatniczych zobowiązań wynika z przyczynowości łączącej się z wyceną opcji. Z jednej strony może to być niezwykle trudne zadanie, ale jednocześnie odpowiednio sformalizowane i w maksymalnym stopniu oddające realia zobowiązań i warunków rynkowych dość łatwe w implementacji kalkulatoryjnej.

Przykładowy wariant wyceny 1

Na (B, S) -rynku prowadzi działalność emitent określonego papieru wartościowego w celu kupna, sprzedaży itp. Aby zostać posiadaczem takiego papieru należy najpierw zapłacić emitentowi określona premię C . Przy tym nabywca ma prawo przedstawienia danego papieru do wykupu w momencie N i otrzymania wypłaty w wysokości f . Taki pochodny papier wartościowy jest znany jako opcja europejskiego typu (na zakup, sprzedaż itd. aktywu), a sama transakcja – kontraktem opcyjnym.

Oczywiste jest, że bardzo ważna jest tu kwestia oceny wartości sprzedaży i kupna opcji oraz sekurytyzowania (hedgingu) płatniczego zobowiązania względem danej opcji. Przede wszystkim należy najpierw sformalizować określenia tych obiektów.

Przyjmujemy, że na (B, S) -rynku (1) zadana jest początkowa wartość kapitału $x > 0$ i płatnicze zobowiązanie (f, N) . Samofinansujący portfel π nazywa się (x, f, N) -hedgingiem (zabezpieczeniem), jeśli kapitał X^π ma własności:

$$X_0^\pi = x, \text{ oraz dla dowolnego } \omega \in \Omega, \quad X_N^\pi \geq f(\omega) \quad (7)$$

Hedging nazywa się minimalnym, jeśli w (7) osiągnięta jest równość. W takim przypadku stwierdza się osiągnięcie płatniczego zobowiązania (f, N) .

Oznaczając przez $\Pi(x, f, N)$ zbiór wszystkich (x, f, N) -hedgingów (zabezpieczeń), przyjmuje się określenie inwestycyjnego kosztu płatniczego zobowiązania (f, N) :

$$C(N) = \inf\{x > 0: \Pi(x, f, N) \neq \emptyset\} \quad (8)$$

Inwestycyjny koszt jest ograniczony, gdyż Ω jest zbiorem skończonym.

Wielkość $C(N)$ nazywa się sprawiedliwą ceną opcyjną. Wynika to stąd, że (f, N) jako płatnicze zobowiązanie opcyjne na zakup, sprzedaż itd. niektórych aktywów, poprzez formułę (8) realizuje zasadę zadowolenia zarówno sprzedawcy, jak i kupującego. Jest tak, ponieważ sprzedawca może na danym rynku „osiągnąć” zobowiązanie (f, N) , a kupujący płaci, w określonym sensie, minimalną premię sprzedawcy.

Wnioski z wprowadzonych założeń urealnienia kierunków modelu rynku

Dla (B, S) -rynku (1) z zadaniem płatniczym zobowiązaniem (f, N) przy opcji europejskiego rodzaju zachowują swoje znaczenie określenia (7) i (8). Amerykańskich opcji, jako skonstruowanych odmiennie od europejskich i wymagających nieco innego ujęcia nie rozpatrujemy w tym opracowaniu, chociaż i dla nich zachowują moc odpowiednie określenia dla kapitału początkowego i kapitału odpowiadającego minimalnemu hedgingowi. W obu przypadkach klasę SF zastępuje klasa GF [8]. Odpowiednie ceny i hedgingi są przy tym nazywane G -cenami i G -hedgingami.

W warunkach zupełnego (B, S) -rynku, przy jedynej martyngałowej mierze P^* i założonym ciągu stóp zwrotu $r_n > -1$ (także amerykańskiego), zadanie wyceny i zabezpieczenia opcji ma adekwatne rozwiązanie w klasie SF .

W podobny sposób analogiczne zadanie można rozpatryć dla klasy G -samofinansujących się strategii (dla uściślenia, strategii z refinansowaniem lub inwestowaniem). Istotne jest przy tym, że można wyjaśnić wtedy, o jaką wielkość różni się sprawiedliwa cena od G -ceny.

Przy przyjętych wyżej założeniach odnośnie do zupełnego rynku zachodzą bowiem trzy własności:

1. Sprawiedliwa cena opcyjna wyraża się w formule:

$$C(N, G) = E^*(\mathcal{E}_N^{-1})(U) \cdot f + \sum_{k=1}^N \mathcal{E}_{k-1}^{-1}(U) \Delta G_k \quad (9)$$

2. Istnieje minimalny (C, f, N) -hedging:

$$\pi^* = ((\beta_n^*, \gamma_n^*))_{n \leq N}$$

określony wzorami:

$$\gamma_n^* = \frac{\alpha_n^* \cdot B_n}{S_{n-1}}$$

$$\beta_n^* = \frac{X_{n-1}^{\pi^*} - \gamma_n^* S_{n-1} - \Delta G_n}{B_{n-1}}$$

gdzie α_T^* – z rozwinięcia [8]:

$$M_n^* = \frac{x}{B_0} + \sum_{k=1}^n \alpha_k^* (\rho_k - r_k), \quad n \leq N$$

a $M_0^* = \frac{x}{B_0}$.

3. Istnieje ściśle określony kapitał minimalnego G -hedgingu.

Rozpatrzony model dotyczył rynku bez arbitrażu charakteryzującego się własnością zupełności. Wiąże się z tym jednoznaczność martyngałowej miary. Względem tej miary prowadzone były wszystkie finansowe kalkulacje i formalne wywody. Jeśli rozważa się niezupełne rynki można także szacować opcyjnie zabezpieczenia oraz minimalne ryzyko. Martyngałowa miara nie jest wtedy jednak jedyna. W tym przypadku należy zredefiniować pojęcie „ceny opcji” (europejskiego rodzaju) ze zobowiązaniem płatniczym w ramach niezupełnego modelu rynku (1). Uczestnik rynku może w tym przypadku występować w charakterze sprzedawcy i w charakterze kupującego opcje. Różne postrzeganie cen przez sprzedawcę i przez kupującego prowadzi w tym przypadku do wyrażenia, ogólnie określając, różnych cen sprzedaży $C^*(N)$ i zakupu $C_*(N)$ i do pojawienia niezerowej różnicy między tymi cenami określanymi powszechnie w literaturze jako *spread*. Przypadek ten jest bardziej złożony merytorycznie i formalnie należy rozpatrywać go inaczej niż w przypadku ujęcia zaprezentowanego w niniejszym opracowaniu.

4. Chaos deterministyczny – schemat systemu

Ujmując zagadnienie generujące chaotyczne warunki kształtowania się cen akcji w warunkach pełnego determinizmu można wprowadzić pojęcie „nieliniowego chaotycznego modelu”. Zbadanie takiego efektu umożliwia ocenę straty na efektywności systemu i przejście systemu w stan chaosu. Konieczna przy tym jest znajomość jednego z istniejących podejść w rozróżnieniu „stochastyczności” i „chaotyczności”. Przedstawione teraz rozróżnienie analizowane jest z zastosowaniem korelacyjnego wymiaru badanych ciągów wielkości, których bliżej nie będziemy omawiać. Umożliwia to wtedy opis metody obliczania górnej i dolnej ceny hedgingowania płatniczego zobowiązania w jednoetapowym modelu rynku dla gwarantowanego przypadku, tj. przy warunku, że stopa procentowa akcji jest „chaotyczną” wielkością.

Chaos i rynki

Ewolucja logarytmicznych stóp zwrotu cen akcji zwykle przedstawiana jest jako ciąg:

$$h = (h_n)$$

gdzie $h_n = \ln \frac{S_n}{S_{n-1}}$ oraz S_n – wartość „ceny” w momencie n , wychodząc z hipotezy, że wielkości te mają stochastyczną naturę, tj.:

$$S_n = S_n(\omega), \quad h_n = h_n(\omega)$$

są wielkościami losowymi, zadanymi na pewnej filtrowanej przestrzeni probabilistycznej $(\Omega, \Phi, (\Phi_n)_{n \geq 1}, P)$ i które modelują stochastyczną nieokreśloność stanów „otoczenia”.

Z drugiej strony stwierdzone zostało już dawno, że nawet zupełnie proste nieliniowe systemy deterministyczne, które można zapisać w postaci:

$$x_{n+1} = f(x_n; \lambda), \quad n = 0, 1, \dots \quad (10)$$

lub:

$$x_{n+1} = f(x_n, x_{n-1}, \dots, x_{n-k}; \lambda), \quad (11)$$

gdzie λ – pewien parametr, mogą generować (przy odpowiednich początkowych warunkach), ciągi $(x_0, x_1, \dots)$, których proveniencja jest podobna do stochastycznych ciągów wartości.

Ta okoliczność uzasadnia pytanie, czy niektóre ekonomiczne, w tym finansowe, szeregi nie są w realnym ich postrzeganiu właśnie nie stochastyczne, a chaotyczne, tzn. takie, które niejako wymuszają modelowanie ich przez deterministyczne nieliniowe systemy. Mogą one prowadzić do efektów obserwowanych przy stochastycznej analizie finansowych danych. Szczególnie znajduje to uzasadnienie w ostatnim okresie, gdy z jednej strony stwierdzono eksperymentalnie reakcję rynków na zachowanie się decydentów politycznych, a z drugiej obserwowane są mało uzasadnione stochastycznymi szokami perturbacje na rynkach finansowych niedające się łatwo wyprofilować metodycznie przez racjonalne działania i nieodpowiadające na łączne losowe reakcje uczestników rynku.

Przykłady nieliniowych chaotycznych systemów

Przytaczając pewne przykłady nieliniowych chaotycznych systemów prezentujemy ich zachowanie się, a także umożliwimy uzasadnienie pytania pojawiającego się w naturalny sposób w takich sytuacjach, a mianowicie, jak określić, czy realizowany dany szereg generowany jest przez stochastyczny czy chaotyczny system.

W aspekcie prognozy przyszłego ruchu cen znacząco ważna jest kwestia, w jakim zakresie można prognozować na podstawie nieliniowych chaotycznych systemów. Okazuje się, że sytuacja nie jest zbyt optymistyczna, gdyż chaotyczne systemy charakteryzuje, niezależnie od ich deterministyczności, duża zmienność ich trajektorii, która może się pojawiać przy niedokładnych danych początkowych, a ponadto zależy istotnie od wartości parametru λ .

Logistyczne przekształcenie

W logistycznej aplikacji przekształceń mającej w ekonomii wiele odniesień rozpatrzmy przekształcenie logistyczne [7]:

$$x \rightarrow Tx \equiv \lambda x(1 - x)$$

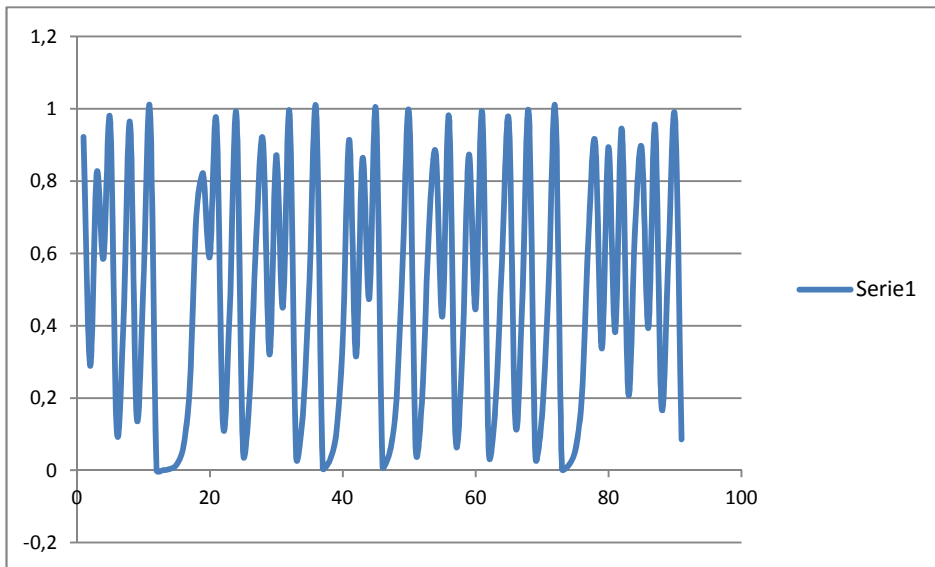
i wywołany przez nie (jednowymiarowy) nieliniowy dynamiczny system:

$$x_n = \lambda x_{n-1}(1 - x_{n-1}), \quad n = 0, 1, \dots, \quad 0 < x_0 < 1 \quad (12)$$

Dla wartości $\lambda \leq 1$ rozwiązania $x_n = x_n(\lambda)$ maleją i są zbieżne do zera przy $n \rightarrow \infty$ i wszystkich $0 < x_0 < 1$. W takim przypadku stan $x_\infty = 0$ można rozpatrywać, jak ten jednoznaczny stabilny stan, do którego zbieżne są wszystkie wartości x_n przy $n \rightarrow \infty$. Przy $\lambda = 2$ wartości x_n są rosnące do 0,5. Zatem w tym przypadku także istnieje jednoznaczne stabilne rozwiązanie $x_\infty = 0,5$, które „przyciąga” wartości x_n przy $n \rightarrow \infty$.

Zwiększając wartość parametru λ łatwo stwierdzić, że w systemie (3), przy $\lambda < 3$ tak jak wcześniej istnieje tylko jedno stabilne rozwiązanie, jednak już przy $\lambda = 3$ powstaje jakościowo nowy efekt, a mianowicie w miarę wzrostu n występują dwa stany stabilności x_∞ , w których na przemian znajduje się system. Taki sam charakter zachowuje system przy zwiększaniu wartości parametru λ , ale system zachowuje się nagle inaczej przy niewielkim wzroście parametru λ , i przy $\lambda = 3,5644...$ takich stanów jest 16, przy $\lambda = 3,5696...$ jest ich już 64, a przy $\lambda = 3,6$ liczba takich stanów jest już nieograniczenie duża. Ten ostatni przypadek tłumaczy się utratą stabilności przez system i przejście systemu w stan chaosu.

Dla $\lambda = 4$ mamy sytuację zbliżoną do losowości probabilistycznej (rys. 1).



Rys. 1. Przypadek $\lambda = 4$, $x_0 = 0,9$

Nieograniczona liczba stanów wyjaśniana jest także w tym przypadku przez utratę stabilności systemu i przejście systemu w stan chaosu, przy tym w pełni znika periodyczny charakter zmiany stanów i system rozpoczyna wykonywanie błędzenia po nieskończonej liczbie stanów. Ważne jest spostrzeżenie, że chociaż system pozostaje deterministyczny, praktycznie nie można przewidzieć, gdzie znajdzie się po pewnym czasie, ponieważ ograniczona dokładność określenia wartości x_n i λ może silnie wpływać na wartości prognozowanych wielkości.

Nie pozostawia zatem wątpliwości fakt, że wartości (λ_k) parametru λ , gdzie zachodzi „rozgałęzienie”, stają się wszystkie „bliżej i bliżej”.

M. Feigenbaum sformułował hipotezę, a O. Lanford wykazał, że (dla wszystkich parabolicznych systemów):

$$\frac{\lambda_k - \lambda_{k-1}}{\lambda_{k+1} - \lambda_k} \rightarrow F, \quad k \rightarrow \infty$$

gdzie $F = 4,669201 \dots$ – stała uniwersalna, nazywana liczbą Feigenbauma.

Parametr $\lambda = 4$ ma w równaniu (12) szczególną rolę – właśnie przy tej wartości ciąg obserwacji odpowiadających (chaotycznych) ciągowi (x_n) przypomina realizację stochastycznego ciągu typu „białego szumu”. W rzeczywistości, jeśli weźmiemy $x_0 = 0,1$ i obliczymy rekurencyjną formułą (12) $x_1, x_2, \dots, x_{1000}$, to empiryczne wartości średniej i odchylenia standardowego wynoszą od-

powiednio, 0,48887 i 0,35742 (z dokładnością do 5 cyfr). Natomiast dla 100 powtórzeń wyniki są następujące: empiryczne wartości średniej i odchylenia standardowego odpowiednio wynoszą: 0,474916 i 0,361261 (dokładnością do 6 cyfr).

Tabela 1

Wartości (empirycznej) korelacyjnej funkcji $\hat{p}(k)$, obliczone dla wartości $x_1, x_2, \dots, x_{1000}$

1	-0,033	11	-0,046	21	-0,008	31	0,038
2	-0,058	12	0,002	22	0,009	32	-0,017
3	-0,025	13	-0,011	23	-0,039	33	0,014
4	-0,035	14	0,040	24	-0,020	34	0,001
5	-0,012	15	0,014	25	-0,008	35	0,017
6	-0,032	16	-0,023	26	0,017	36	-0,052
7	-0,048	17	-0,030	27	0,006	37	0,004
8	0,027	18	0,037	28	-0,004	38	0,053
9	-0,20	19	0,078	29	-0,019	39	-0,021
10	-0,013	20	0,017	30	-0,076	40	0,007

Źródło: Opracowanie własne na podstawie [6].

Z podanych w tabeli wartości funkcji korelacyjnej jest widoczne, że wielkości (x_n) , wywołane przez logistyczne przekształcenie z $\lambda = 4$ praktycznie można uważać jako nieskorelowane i w tym znaczeniu ciąg (x_n) może być nazwany „chaotycznym białym szumem”.

Interesująca jest uwaga, że dla systemu $x_n = 4x_{n-1}(1 - x_{n-1})$, $n = 1, \dots$, $0 < x_0 < 1$, istnieje niezmienniczy rozkład P , tj. taki, że $P(T^{-1}A) = P(A)$ dla dowolnego borelowskiego zbioru A z przedziału $(0, 1)$, którego gęstość:

$$p(x) = \frac{1}{\pi[x(1-x)]^{1/2}}, \quad x \in (0, 1) \quad (13)$$

Wynika z tego, że jeśli przyjąć losową początkową wartość x_0 z gęstością rozkładu prawdopodobieństwa $p = p(x)$, to losowe wielkości x_n , $n \geq 1$ będą z tego samego rozkładu, z którego pochodzi x_0 .

Należy tu stwierdzić, co wynika z teorii probabilistyki, że w uzyskanym w ten sposób stochastycznym systemie (x_n) cała „losowość” jest całkowicie określona przez początkową wartość x_0 , a dynamika przejść $x_n \rightarrow x_{n+1}$ zadana jest w deterministyczny sposób w relacji (12).

Przy rozkładzie zadanym funkcją gęstości (4) nietrudno jest stwierdzić, że wartość oczekiwana $Ex_0 = 0,5$, $Ex_0^2 = \frac{3}{8}$, $D^2x_0 = \frac{1}{8}$ ($D^2 = (0,35355^2)$) (średnia z wartością 0,48887 i odchylenie standardowe 0,35742, podane wyżej) i:

$$p(k) \equiv \frac{Ex_0x_k - Ex_0E_k}{\sqrt{Dx_0Dx_k}} = \begin{cases} 1, & \text{jeśli } k = 0 \\ 0, & \text{jeśli } k \neq 0 \end{cases}$$

Przykładowe przekształcenia w modelowaniu finansowych wskaźników w okresach kryzysu finansowego

1. Przekształcenie Bernoulliego

$$x_n = 2x_{n-1} \pmod{1}, n = 1, 2, \dots, x_0 \in (0, 1)$$

W tym przypadku niezmienniczy jest jednostajny rozkład z gęstością $p(x) = 1$, $x \in (0, 1)$. Podstawowe charakterystyki w tym rozkładzie dla wielkości losowej x_0 wynoszą:

$$Ex_0 = \frac{1}{2}, Ex_0^2 = \frac{1}{3}, Dx_0 = \frac{1}{12}, p(k) = 2^{-k}, k = 0, 1, \dots$$

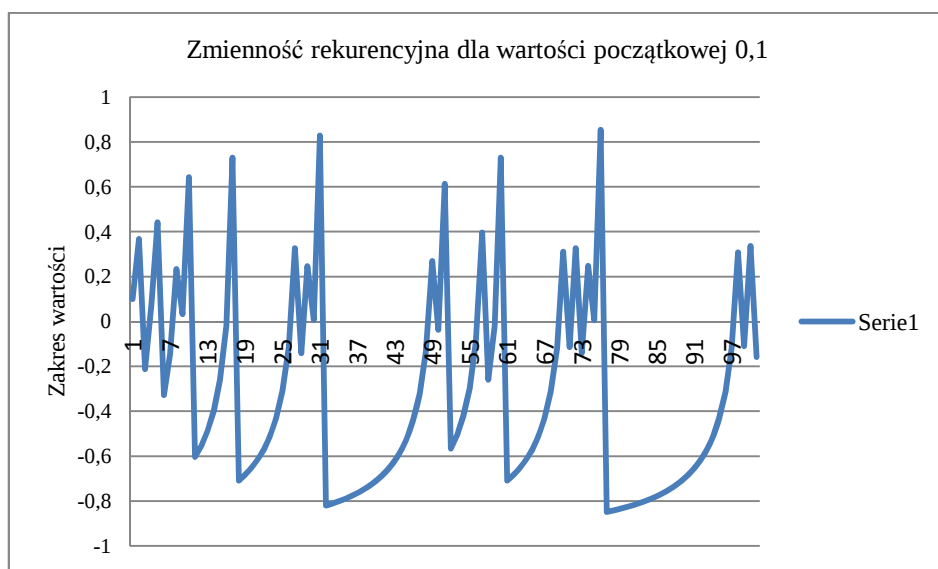
2. Przekształcenie „namiotowe”

$$x_n = 1 - |1 - 2x_{n-1}|, n = 1, 2, \dots, x_0 \in (0, 1)$$

Tu także, jak dla przekształcenia Bernoulliego, niezmienniczy jest rozkład jednostajny w przedziale $(0, 1)$. Ponadto $Ex_0 = \frac{1}{2}$, $Ex_0^2 = \frac{1}{3}$, $Dx_0 = \frac{1}{12}$, $p(k) = 0$, $k \neq 0$.

3. Przekształcenie pierwiastkowe

$$x_n = 1 - 2\sqrt{x_{n-1}}, n = 1, \dots, x_0 = (-1, 1)$$



Rys. 2. Wykres zmian wartości rekurencyjnych w przekształceniu pierwiastkowym

Niezmienniczy jest tutaj rozkład jednostajny na odcinku $(-1, 1)$ z gęstością $p(x) = \frac{1-x}{2}$, przy tym $Ex_0 = -\frac{1}{3}$, Ex_0^2 , $Dx_0 = \frac{2}{9}$.

Wskazane przykłady nieliniowych dynamicznych systemów są istotne w różnych aspektach. Po pierwsze, można zauważyć, na przykładzie logistycznego systemu, którego rozwój jest „binarny”, że wyrażenie wyraża się idea chaotyczności. Po drugie, kształtowanie się takich systemów, scharakteryzowanych własnością chaotyczności, przytacza na myśl ich zastosowanie przy konstrukcji modeli ewolucji finansowych indeksów, a szczególnie w okresach kryzysowych. Dla takich okresów właściwa jest właśnie „chaotyczność”, a nie „stochastyczność”.

Okoliczność, że formalnie deterministyczne systemy mogą przejawiać właściwości typu „stochastycznego białego szumu” jest znana od dawna i nie jest czymś nieoczekiwanym. Dlatego powstają dwie kwestie odnoszące się do tego, jak rozróżniać „stochastyczne” i „chaotyczne systemy” oraz czy można w zasadniczy sposób zdecydować, jaka jest istotna natura „nieregularności” finansowych danych – „stochastyczna” czy „chaotyczna”.

Przedstawimy ujęcie mające główne znaczenie przy rozróżnianiu „stochastyczności” i „chaotyczności” funkcji:

$$C(\varepsilon) = \lim_{N \rightarrow \infty} \frac{\emptyset(N, \varepsilon)}{N^2} \quad (14)$$

gdzie $\emptyset(N, \varepsilon)$ – liczba takich par (i, j) , $i, j \leq N$, dla których w rozpatrywanym ciągu (x_n) elementy tego ciągu są odległe o mniej niż ε , tzn.:

$$|x_i - x_j| < \varepsilon$$

Oprócz funkcji $C(\varepsilon)$ rozważa się funkcję:

$$C_m(\varepsilon) = \lim_{N \rightarrow \infty} \frac{\emptyset_m(N, \varepsilon)}{N^2}$$

gdzie $\emptyset(N, \varepsilon)$ – liczba takich par (i, j) , dla których wszystkie współrzędne wektorów $(x_i, \dots, x_{i+m-1})$ i $(x_j, \dots, x_{j+m-1})$ dla $i, j \leq N$ różnią się nie więcej niż ε . W wypadku $m = 1$ mamy $\emptyset_1(N, \varepsilon) = \emptyset(N, \varepsilon)$.

Dla stochastycznych ciągów (x_n) mających cechy „białego szumu” przy małych wartościach ε funkcja spełnia relację:

$$C_m(\varepsilon) \sim \varepsilon^{\nu_m} \quad (15)$$

gdzie „fraktalny” wykładnik $\nu_m = m$. Własnościom w rodzaju (15) czyni zadość także i wiele deterministycznych systemów, w tym logistyczny system (12). Wykładnik ν_m nazywany jest także korelacyjnym wymiarem.

Idea rozróżniania „stochastycznych” i „chaotycznych” ciągów wychodzi z takiej obserwacji, że korelacyjny wymiar w takich ciągach jest różny. W ciągach „stochastycznych” jest większy niż w „chaotycznych”.

Oceny korelacyjnego wymiaru

W charakterze ocen korelacyjnego wymiaru ν_m naturalne jest rozpatrzenie wielkości:

$$\tilde{\nu}_{m,j} = \frac{\ln C_m(\varepsilon_j) - \ln C_m(\varepsilon_{j+1})}{\ln \varepsilon_j - \ln \varepsilon_{j+1}}$$

lub:

$$\tilde{\nu}_m(k) = \frac{1}{k} \sum_{j=1}^k \tilde{\nu}_{m,j}$$

gdzie $\varepsilon_j = \varphi^j$, $0 < \varphi < 1$.

Z wyników znanych z literatury [8] wynika po pierwsze, jednorodność „fraktalnej” struktury „korelacyjnego wymiaru” indeksów IBM i S & P500, po drugie, że dla tych indeksów ciągi (h_n) z $h_n = \ln \frac{S_n}{S_{n-1}}, n \geq 1$ zmierzają do siebie szybciej niż stochastyczny „biały szum”. Spostrzeżenie to nie jest podstawą do odrzucenia hipotezy o tym, że bliskimi właściwościami mogą charakteryzować się także inne „chaotyczne” ciągi z wielkim „korelacyjnym rozmiarem”.

W praktyce analizowania problemu rozróżniania „chaotyczności” i „stochastyczności” znajduje także zastosowanie podejście, w którym rozpatrywane jest kształtowanie się rozkładu prawdopodobieństwa systemu.

5. Wariant rozwiązania zadania rekonstrukcji operatora ewolucji dla rynków futures

W zastosowaniach rozpatrywane są różne możliwe rozwiązania zadania rekonstrukcji operatora ewolucji dla rynków futures. Podstawową przesłanką jest to, że dowolny wybór nieliniowości bez wprowadzenia apriorycznej informacji lub specjalnego uprzedniego badania obiektu nie zawsze umożliwia wybór udanej rekonstrukcji. Dlatego na danym etapie modelowania szerokie zastosowanie ma dotąd dla prognozowania rynkowych charakterystyk analiza techniczna. Reguły tej teorii uwzględniają fakt, że w dynamice rynku akcji istotne są trzy podstawowe źródła informacji: ceny akcji, wielkość sprzedaży i otwarte zlecenia. Wielkość obrotów i otwarte zlecenia nie są arbitralnie znaczące, ale mimo tego są ważnymi czynnikami wpływającymi na formowanie cen akcji. Otwarte zlecenia są ilością niezamkniętych pozycji na końcu dziennej sesji.

Tak zobrazowany proces jest podstawą modelu prognozowania cen na rynku futures i powinien opisywać zmiany trzech komponent rynkowych – cena kontraktu, wielkość obrotów, otwarte pozycje. Istniejąca relacja między opisanymi wskaźnikami ekonomicznymi wyrażona jest krzyżującymi się iloczynami odpowiednich fazowych zmiennych występujących w modelu prognozowania cen na rynku futures. Dane parametry są zmiennymi na pewnym dostatecznie dużym odcinku czasu, ale kawałkami stałymi na niewielkim badanym przedziale czasu-kroku prognozy. Taki model, analizowany na podstawie teorii deterministycznego chaosu, wskazuje, że wiele losowych ekonomicznych zjawisk jest bardziej przewidywalnych niż przyjęto sądzić.

Literatura

1. Andritzky B., *Sovereign Default Risk Valuation. Implication of Debt Crises and Bond Restructurings*, Verlag, Berlin 2006.
2. Hull J., *Futures and Other Deivative Securities*, Englewood Cliffs, Prentice-Hall 1992.
3. Karatzas J., Shreve S.E., *Methods of mathematical finance*, Springer Verlag, New York 1999.
4. Lipcer N., Szirajew A.R., *Statystyka procesów stochastycznych*, PWN, Warszawa 1981.
5. Mandelbrot B., *Fractals and scaling in finance: discontinuity, concentration, risk*, Springer 1997.
6. *Podstawy stochastycznej finansowej matematyki*, T. 1. Fakty. Modele, T. 2. Teoria, FAZIS, Moskwa 1998.
7. Szirajew I., *Opcje i ryzyko, prawdopodobieństwo, zabezpieczenia i chaos. Matematyka finansów*, URSS, Moskwa 1999.
8. Szirajew W., *Fianansowyie rynki: Neironye eti, chaos, nichinejnaia dinamica*, Dom Książki Librocom, Moskwa 2009.
9. Wilmott P., Howison S., Dewynne J., *The Mathematics of Financial Derivatives*, Cambridge University Press 1997.

CHAOTIC RESPONSE OF THE FINANCIAL MARKETS – PROBABILISTIC ASPECTS OF PAYMENT SECURITY VALUATION AND CAPITAL MARKET

Summary

Considered in developing the financial model of the exemplification of the market refers to the complete markets. Developed the idea not-self-financing the strategy at a fair valuation of the possible options. The appropriate development of this theme is the introduction to the issue of the financial model of incomplete markets and to structure the equity portfolio under the assumption of statistical indeterminacy. The derived formulas in the article is a basic introduction to the analysis of the financial market, but the aspect of perspective on this subject with seemingly very formalized, leading to appraise the relevant hedging approach equivalent security.

The following article about the chaotic financial data responsive to capital markets. Examined aspect of distinguishing chaotic and stochastic defined in terms of looking at this problem. Discusses the correlation dimension as an assessment of the degree of chaos in time series data. Attention has been returned to the issue in various applications possible solutions to the tasks for the reconstruction of the evolution operator of futures markets. The basic premise is that any choice of non-linearities without introducing a priori information or special prior studies do not always object to select the successful reconstruction.

Konstancja Poradowska

Uniwersytet Ekonomiczny we Wrocławiu

MODELE SUBIEKTYWNE W KONSTRUKCJI PROGNOZ DŁUGOOKRESOWYCH

Wprowadzenie

Dynamiczny rozwój gospodarki, cywilizacji i postępu technologicznego stwarza potrzebę modelowania i prognozowania nowych zjawisk, czego potwierdzeniem może być wciąż wzrastająca w Polsce i na świecie popularność badań typu *foresight*. Główną przyczyną trudności bywa tu jednak brak dostatecznej liczby danych empirycznych, pozwalających na „klasyczną budowę” matematycznego modelu rzeczywistości. Rutynowym podejściem jest w takiej sytuacji wykorzystanie heurystycznych metod prognozowania, opartych na opiniach ekspertów, które mogą być zebrane np. za pomocą ankiety delfickiej. Badania pokazują jednak, że trafność prognoz formułowanych bezpośrednio przez ekspertów rzadko bywa zadowalająca, zwłaszcza w zestawieniu z prognozami otrzymanymi na podstawie formalnego modelu prognostycznego [12]. Trudności te nasilają się, gdy np. na potrzeby długookresowych scenariuszy rozwoju wymagana jest konstrukcja całej trajektorii prognoz, sięgającej wielu okresów naprzód – w przypadku badań *foresight* nawet kilkudziesięciu lat. Alternatywą dla „tradycyjnych” metod heurystycznych może być wówczas budowa tzw. formalnego modelu subiektywnego (modelu formalnego II rodzaju), którego parametry ocenia się na podstawie subiektywnej informacji pozyskanej od ekspertów. W zależności od zakresu posiadanej informacji może to być model przyczynowo-skutkowy [6; 7] lub model tendencji rozwojowej [4; 5; 14].

Wybrane aspekty budowy i praktycznego wykorzystania subiektywnych modeli prognostycznych stanowią podstawowy przedmiot rozważań zamieszczonych w niniejszym opracowaniu. Celem głównym jest wskazanie przydatności takich modeli w konstrukcji długookresowych prognoz i scenariuszy roz-

woju nowych technologii. Rozważania teoretyczne zostaną uzupełnione o realne przykłady analizy danych, pozyskanych w badaniu *foresight* „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”, realizowanego przez Główny Instytut Górnictwa w Katowicach.

1. Subiektywne i obiektywne modele prognostyczne

Jedną z klasyfikacji metod prognozowania jest ich podział na metody ilościowe i jakościowe. Metody ilościowe są oparte na formalnych modelach prognostycznych (np. na modelach ekonometrycznych), zbudowanych na podstawie obiektywnych danych o kształtowaniu się zmiennej prognozowanej i zmiennych objaśniających w przeszłości. Przedstawienie zależności pomiędzy poszczególnymi zmiennymi w postaci matematycznego równania umożliwia rozważenie różnych scenariuszy rozwoju przyszłości. Takie modele uwzględniają jednak wyłącznie prawidłowości występujące w danych prognostycznych, stąd pozwalają osiągnąć dobre rezultaty, jeżeli w horyzoncie prognozy nie zajdą istotne zmiany w czynnikach wpływających na prognozowane zjawisko i w sposobie ich oddziaływania, a więc głównie w przypadku prognozowania krótkookresowego. Zdarzenia, które nie zostały zaobserwowane w przeszłości, lecz są oczekiwane w okresie prognozy mogą być uwzględnione poprzez zastosowanie jakościowych metod prognozowania. Metody jakościowe są oparte na subiektywnych sądach eksperckich, czyli na modelach myślowych (nieformalnych), których nie da się przedstawić w sformalizowanym języku matematyki. Praktyka pokazuje, że eksperci bywają często zbyt optymistami, dlatego prognozy powstałe wyłącznie na podstawie modeli myślowych mogą wykazywać tendencję do obciążoności* [1; 3; 8].

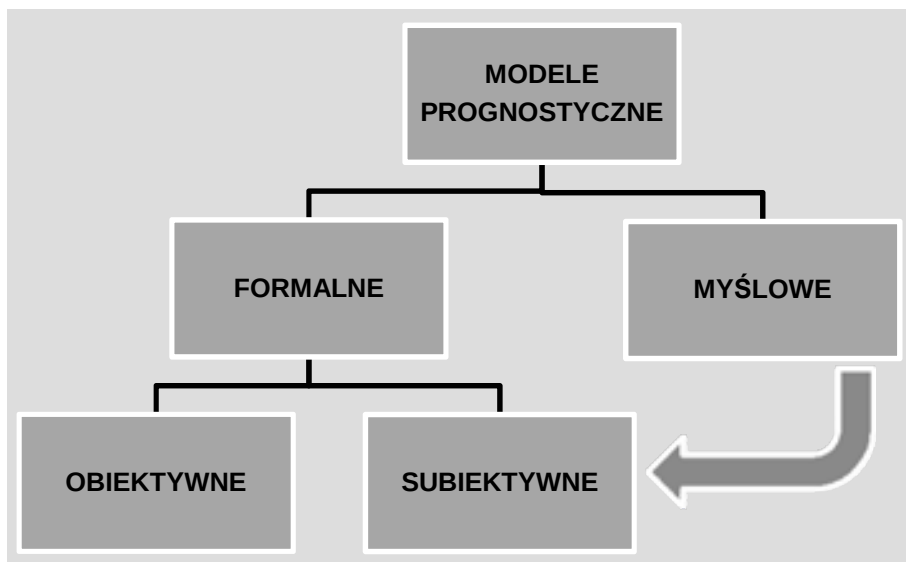
Rozważając wady i zalety obu rodzajów metod prognostycznych można dojść do wniosku, że aby przy formułowaniu prognozy uwzględnić wszystkie dostępne informacje zachodzi potrzeba integracji ilościowych i jakościowych metod prognozowania. Do procedur takiej integracji (obok kombinacji prognoz oraz ich korygowania [4, s. 190-191]) należy prognozowanie na podstawie subiektywnych modeli formalnych (modeli formalnych II rodzaju). Wartości parametrów takich modeli, w przeciwieństwie do powszechnie stosowanych obiektywnych modeli formalnych, nie są szacowane klasycznymi metodami statys-

* To znaczy błędy wyznaczonych przez eksperta prognoz bywają jednokierunkowe – prognozy są systematycznie przeszacowywane lub niedoszacowywane.

tycznymi, lecz określone na podstawie ocen ekspertów, a zatem z użyciem modeli myślowych. Klasyfikację modeli prognostycznych przyjętą w prezentowanym opracowaniu przedstawiono na rys. 1.

Modele subiektywne w konstruowaniu prognoz są szczególnie użyteczne, gdy:

- sądy ekspertów wskazują, że zaobserwowane dotychczas prawidłowości mogą zaniknąć w przyszłości,
- prognosta nie dysponuje danymi pozwalającymi na szacowanie parametrów modelu metodami statystycznymi, np. gdy prognoza dotyczy zjawiska nowego.



Rys. 1. Schemat klasyfikacji modeli prognostycznych

2. Subiektywne modele tendencji rozwojowej

Znane z literatury przedmiotu subiektywne modele tendencji rozwojowej służą do opisu dynamiki sprzedaży nowych produktów [4; 5; 14]. Prognosta przyjmuje założenie o postaci funkcyjnej modelu w oparciu o spodziewany kształt krzywej życia produktu. Wykorzystywane są w tym celu następujące funkcje:

1) liniowa:

$$Y_t = \alpha + \beta t \quad (1)$$

2) wykładnicza:

$$Y_t = \alpha(1 + g)^t \quad (2)$$

oraz, jeśli dodatkowo przyjmuje się założenie o skończonym potencjale rynku:

3) wykładnicza odwrotnościowa (z asymptotą poziomą):

$$Y_t = \alpha - \beta g^t, \quad g < 1 \quad (3)$$

4) logistyczna:

$$Y_t = \frac{1}{\alpha - \beta g^t} \quad (4)$$

gdzie:

t – zmienna czasowa,

α, β, g – parametry modelu.

Oceny parametrów wyznacza się na podstawie sądów eksperta lub grupy ekspertów, które dotyczą:

- w przypadku funkcji liniowej i wykładniczej – wartości dwóch zmiennych losowych: wielkości sprzedaży w pierwszym okresie istnienia produktu na rynku (Y_1) oraz wielkości sprzedaży w jednym z późniejszych okresów, po ustabilizowaniu się (Y_n),
- w przypadku funkcji wykładniczej odwrotnościowej i logistycznej – wartości trzech zmiennych losowych: wielkości sprzedaży w pierwszym okresie istnienia produktu na rynku (Y_1), wielkości sprzedaży w jednym z późniejszych okresów (Y_n) oraz poziomu nasycenia rynku (Y_∞).

Formuły pozwalające na wyznaczenie parametrów α, β, g wraz z wykreśami odpowiednich funkcji (1)-(2) przedstawiono w tab. 1. Prognozę y_T^* na dowolny okres $T > 1$ wyznacza się poprzez ekstrapolację zbudowanego modelu.

Tabela 1

Formuły ocen parametrów wybranych subiektywnych modeli tendencji rozwojowej

Postać funkcji trendu	Oceny parametrów modelu		
	α	β	g
Liniowa	$\alpha = y_1 - \beta$	$\beta = \frac{y_n - y_1}{n - 1}$	
Wykładnicza	$\alpha = \frac{y_1}{1 + g}$		$g = \sqrt[n-1]{\frac{y_n}{y_1}} - 1$
Wykładnicza odwrotnościowa	$\alpha = y_\infty$	$\beta = \frac{\alpha - y_1}{g}$	$g = \sqrt[n-1]{-\frac{y_n - \alpha}{\alpha - y_1}}$
Logistyczna	$\alpha = \frac{1}{y_\infty}$	$\beta = \frac{\alpha - \frac{1}{y_1}}{g}$	$g = \sqrt[n-1]{-\frac{\frac{1}{y_n} - \alpha}{\alpha - \frac{1}{y_1}}}$

Szerszą prezentację zagadnienia prognozowania na podstawie subiektywnych modeli tendencji rozwojowej wraz z propozycjami oceny stopnia niepewności prognoz można znaleźć w pracy [10].

3. Wybrane modele dyfuzji innowacji

Pierwszym szeroko rozwiniętym teoretycznie modelem dyfuzji jest zaproponowany przez F.M. Bassa model wzrostu nowego produktu. Model ten stosowano do przewidywania dyfuzji innowacji w handlu detalicznym, technologii przemysłowej, rolnictwie oraz na rynku dóbr trwałego użytku. Bazuje on na założeniu, że istnieje analogia pomiędzy dyfuzją innowacji a rozprzestrzenianiem się epidemii [2]. Model Bassa można opisać za pomocą następującego równania różniczkowego:

$$\frac{dN(t)}{dt} = \left[p + \left(\frac{q}{M} \right) N(t) \right] [M - N(t)] \quad (5)$$

które ma rozwiązanie postaci:

$$N(t) = B(t, M, p, q) = M \frac{1 - e^{-(p+q)t}}{1 + \frac{q}{p} e^{-(p+q)t}} \quad (6)$$

gdzie:

$dN(t)/dt$ – tempo zmian w skumulowanej liczbie nabywców, którzy wdrożyli innowację w czasie t ,

$N(t)$ – ogólna liczba nabywców, którzy wdrożyli innowację w czasie t ,

M – potencjał rynkowy,

p – współczynnik innowacji (prawdopodobieństwo pierwszego zakupu przez grupę innowatorów),

q – współczynnik imitacji.

Pierwszy czynnik modelu (5) reprezentuje prawdopodobieństwo wdrożenia innowacji, drugi – liczbę potencjalnych nabywców, którzy jeszcze tego nie dokonali. W modelu przyjmuje się, że na skłonność do przyjęcia innowacji wpływają dwa podstawowe rodzaje środków komunikacji – masowa oraz ustna. Dzieli się zatem konsumentów na: innowatorów (którzy działają pod wpływem komunikacji masowej) oraz imitatorów (naśladowców, którzy działają pod wpływem komunikacji ustnej). Przy braku danych empirycznych z przeszłości, parametry p i q modelu Bassa można określić następująco [6]:

- na podstawie danych dotyczących produktów o analogicznym cyklu życia,
- przyjąć wartości *a priori*, np. $p = 0,003$, $q = 0,5^*$,
- na podstawie sądów eksperckich (wykorzystując np. uogólnioną metodę najmniejszych kwadratów).

Swego rodzaju rozwinięcie modelu Bassa stanowi model E.M. Rogersa [11], który dodatkowo wyjaśnia strukturę komunikacji pomiędzy grupami innowatorów i imitatorów. W modelu Rogersa zakłada się, że w związku z występowaniem w procesie dyfuzji relacji interpersonalnych krzywa adaptacji ma rozkład normalny. Wykorzystując parametry rozkładu normalnego Rogers skategoryzował konsumentów według tempa przyjmowania innowacji i podzielił ich na 5 grup: innowatorów, wczesnych naśladowców, wczesną większość, późną większość, maruderów [http://www.zie.pg.gda.pl/photo/upd/100111173052_wykreslistonic_large.jpg]. Model można opisać następującym równaniem:

* Przyjęcie takich wartości proponuje F. Bass na założonej przez siebie stronie internetowej o tematyce modeli Bassa [www.bassbasement.org]. Lilien i Rangaswamy przyjmują tu średnią wartość parametrów oszacowanych dla określonej grupy produktów.

$$\frac{dN(t)}{dt} = \frac{a \cdot M \cdot e^{-a(t-b)}}{[1 + e^{-a(t-b)}]^2} \quad (7)$$

którego rozwiązaniem jest krzywa logistyczna:

$$N(t) = L(t, M, a, b) = \frac{M}{1 + e^{-a(t-b)}} \quad (8)$$

gdzie:

$dN(t)/dt$ – tempo zmian w skumulowanej liczbie nabywców, którzy wdrożyli innowację w czasie t ,

$N(t)$ – ogólna liczba nabywców, którzy wdrożyli innowację w czasie t ,

M – potencjał rynkowy,

a, b – parametry modelu.

Zakładając, że rozwój zjawiska będzie się kształtował zgodnie z modelem Rogersa można tak sformułować pytania do ekspertów, aby otrzymać informację o punktach szczególnych modelu (zob. rys. 2), które posłużą do oceny parametrów a i b . W zależności od sytuacji można wybrać jeden spośród następujących zestawów pytań* [15]:

Zestaw I

1. W którym okresie (t^*) rynek innowacji osiągnie połowę potencjału? $\rightarrow b$.
2. Ile nowych jednostek w okresie t^* zaadaptuje innowację? $\rightarrow a$.

Parametry modelu wyznacza się tu z zależności:

$$\max \frac{dN(t)}{dt} = \frac{aM}{4} \quad \text{dla} \quad t = b \quad (9)$$

Zestaw II

1. W którym (możliwie krótkim) przedziale czasowym $[t_1, t_2]$ najwięcej nowych użytkowników wdroży innowację? $\rightarrow b$.
2. Jaka to będzie liczba (n) użytkowników? $\rightarrow a$.

Parametry modelu wyznacza się z zależności:

* W poszczególnych pytaniach po symbolu „ $\rightarrow$ ” podano parametr, którego wartość otrzymuje się w wyniku odpowiedzi.

$$n \approx (t_2 - t_1) \frac{aM}{4}, \quad \frac{t_2 - t_1}{2} = b \quad (10)$$

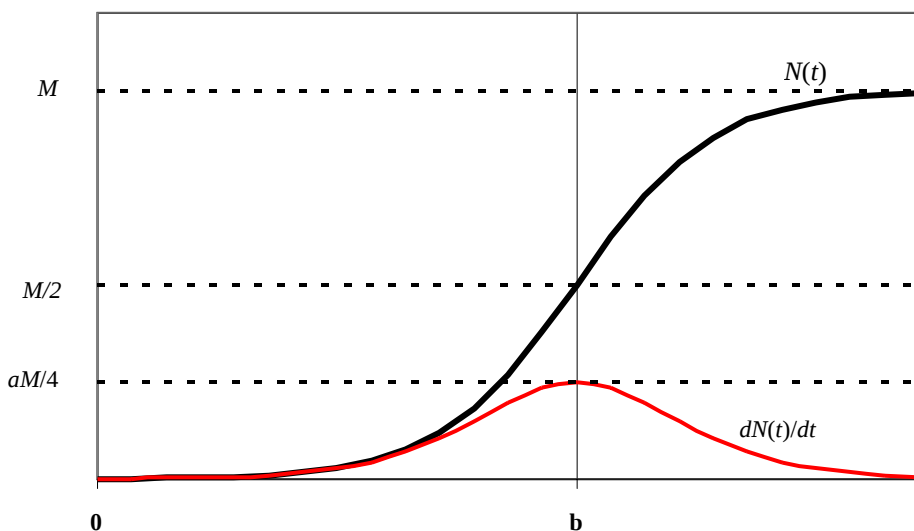
Zestaw III

1. W którym okresie (t_s) zostanie osiągnięte $u \cdot 100\%$ potencjału?
2. Jaki czas jest potrzebny (Δt), licząc od okresu t_s , aby osiągnąć $v \cdot 100\%$ potencjału?

Znając wartości t_s , Δt , u oraz v , parametry a i b wyznacza się ze wzorów:

$$a = \frac{1}{\Delta t} \left[\ln \left(\frac{1}{u} - 1 \right) - \ln \left(\frac{1}{v} - 1 \right) \right], \quad b = t_s + \Delta t \frac{\ln \left(\frac{1}{u} - 1 \right)}{\ln \left(\frac{1}{u} - 1 \right) - \ln \left(\frac{1}{v} - 1 \right)} \quad (11)$$

Należy zauważyć, że model Rogersa pokrywa się z logistycznym modelem tendencji rozwojowej, opisanym równaniem (4), a na podstawie odpowiedzi na zestaw pytań III można również otrzymać wielkości służące do oceny parametrów modelu (zob. tab. 1). Jeżeli prognosta decyduje się na wykorzystanie subiektywnego modelu logistycznego można w zależności od sytuacji wybrać taki sposób oceny parametrów, aby ekspertom najłatwiej było określić wielkości niezbędne do ich wyznaczenia.



Rys. 2. Krzywe Rogersa oraz ich punkty szczególne

4. Subiektywne modele przyczynowo-skutkowe

W przypadku modeli przyczynowo-skutkowych prognosta na wstępie przyjmuje założenie o postaci funkcji $y_t = f(x_t)$ opisującej wpływ zmiennej objaśniającej X na zmienną prognozowaną Y w czasie t . W szczególności może to być funkcja liniowa, wykładnicza, wielomianowa, logarytmiczna, logistyczna [6; 13]. Parametry są określane na podstawie odpowiedzi ekspertów na odpowiednio sformułowane pytania, np.:

1. Jaka jest aktualna/bazowa wartość zmiennych X i Y ?
2. Jakiego poziomu Y należałoby oczekiwać, gdyby wartość X została zredukowana do 0?
3. Jaki (maksymalny) poziom osiągnię Y przy nieograniczonym X ?
4. Ile wyniosłoby Y , gdyby X zwiększono/obniżono o 50%?*

Najlepiej znanym subiektywnym modelem przyczynowo-skutkowym jest tzw. model ADBUDG (Advertising Budget Model), zaproponowany przez Little'a [7] na potrzeby problemu decyzyjnego dotyczącego ustalenia optymalnych wydatków na reklamę.

Zależność wielkości sprzedaży (Y) od wydatków na reklamę (X) została tam opisana funkcją logistyczną jako:

$$y_t = f(x_t) = a + (b - a) \frac{x_t^c}{d + x_t^c} \quad (12)$$

Parametry a i b można otrzymać jako odpowiednie granice funkcji (12) na podstawie odpowiedzi na pytania 2 oraz 3:

$$a = \lim_{x_t \rightarrow 0} f(x_t), \quad b = \lim_{x_t \rightarrow \infty} f(x_t) \quad (13)$$

Parametry c i d są rozwiązaniem układu równań:

$$\begin{cases} a + (b - a) \frac{x_0^c}{d + x_0^c} = y_0 \\ a + (b - a) \frac{(1,5 \cdot x_0)^c}{d + (1,5 \cdot x_0)^c} = y_1 \end{cases} \quad (14)$$

gdzie: x_0, y_0 to wartości bazowe zmiennych X i Y otrzymane w wyniku odpowiedzi na pytanie 1, natomiast y_1 to wartość Y określona w pytaniu 4.

* Zamiast 50% można zapytać o inną wartość, jeżeli w danej sytuacji prognostycznej byłaby ona bliższa intuicji ekspertów.

5. Przydatność modeli w badaniach *foresight* – przykłady

Przedstawione modele dyfuzji zostały wykorzystane do konstrukcji prognoz rozwoju nowych technologii energetycznych na potrzeby badania *foresight*: „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”, prowadzonego w Głównym Instytucie Górnictwa w Katowicach^{*}. Poniżej przedstawiono wybrane wyniki dotyczące rozwoju technologii OZE.

We wcześniejszych etapach badania *foresight* panele ekspertów dostarczyły ocen:

- wielkości produkcji energii z OZE w Polsce w latach 2010, 2020, 2050,
- rynkowego potencjału energetycznego M dla poszczególnych źródeł energii do 2050 r.

Na podstawie tych informacji dla rozwoju poszczególnych technologii OZE zostały wyznaczone wykładnicze odwrotnościowe modele tendencji rozwojowej oraz modele dyfuzji: Rogersa^{**} i Bassa. Opinie ekspertów oraz otrzymane oceny parametrów modeli przedstawiono w tab. 2.

Tabela 2

Opinie ekspertów dotyczące rozwoju technologii OZE oraz wyznaczone na ich podstawie oceny parametrów modeli dyfuzji

Technologia OZE	Opinie ekspertów			Model wykładniczy odwrotnościowy			Model Rogersa		Model Bassa	
	2010 r.	2020 r.	2050 r.	M	β	g	a	b	p	q
1	2	3	4	5	6	7	8	9	10	11
Kolektory słoneczne płaskie/próżniowe	66	2350	4000	5500	5738,53	0,95	0,41	11,71	0,04	0
Fotowoltaika	1,3	450	1000	4000	4046,58	0,99	0,60	14,46	0,01	1E-14
Energetyka wodna klasyczna i szczytowa	2200	2800	10000	12500	10362,00	0,99	0,03	52,25	0,02	0,03
Energetyka wiatrowa wielkiej skali	1400	14000	22000	25000	25472,05	0,93	0,31	10,21	0,07	1E-09
Pompy ciepła i geotermia	320	2700	25000	50000	24931,54	0,99	0,22	24,12	0,01	0,12

* Nr POIG.01.01.01-00-007/08.

** Wartości teoretyczne otrzymane na podstawie modelu Rogersa pokrywają się z wartościami teoretycznymi logistycznego modelu tendencji rozwojowej (4).

cd. tabeli 2

1	2	3	4	5	6	7	8	9	10	11
Mikro-generacja na bazie biomasy	0,001	10	25	50	51,13	0,98	0,94	12,47	0,02	1E-11
Mikro-energetyka wiatrowa	0,001	1	10	20	20,10	0,99	0,70	15,23	0,01	1E-08

Wielkości dla mikrogeneracji na bazie biomasy zostały określone w MWh, dla pozostałych technologii w GWh.

Parametry modelu wykładniczego odwrotnościowego wyznaczono na podstawie formuł zawartych w tab. 1. Jako y_1 przyjęto wielkości produkcji energii określone przez ekspertów dla 2010 r., jako y_n ($n = 11$) wielkości określone dla 2020 r., natomiast y_∞ to odpowiednie potencjały M .

Do oceny parametrów modelu Rogersa wykorzystano zestaw pytań III. Wartości określone przez ekspertów zostały tak przeliczone, aby za okres t_s , występujący w pierwszym pytaniu przyjęto 2010 r. Następnie dla każdej technologii obliczono, jaką część potencjału rynkowego stanowi wartość określona dla 2010 r., otrzymując w ten sposób wartość u . Podobnie postąpiono z wartością dla 2020 r., otrzymując w ten sposób wartość v oraz przedział Δt , wynoszący 10 lat. Parametry a i b obliczono z formuł (11).

Parametry modelu Bassa oceniono na podstawie wszystkich czterech wartości określonych przez ekspertów dla poszczególnych technologii. Wstępnie przyjęto $p = 0,003$, $q = 0,5$. Po wyznaczeniu wartości teoretycznych na lata 2010-2050 wielkości p i q zostały tak skorygowane, aby zminimalizować średnią ważoną kwadratów odchyleń:

$$\alpha = 0,5 \cdot (y_{2010} - \hat{y}_{2010})^2 + 0,3 \cdot (y_{2020} - \hat{y}_{2020})^2 + 0,2 \cdot (y_{2050} - \hat{y}_{2050})^2 \quad (15)$$

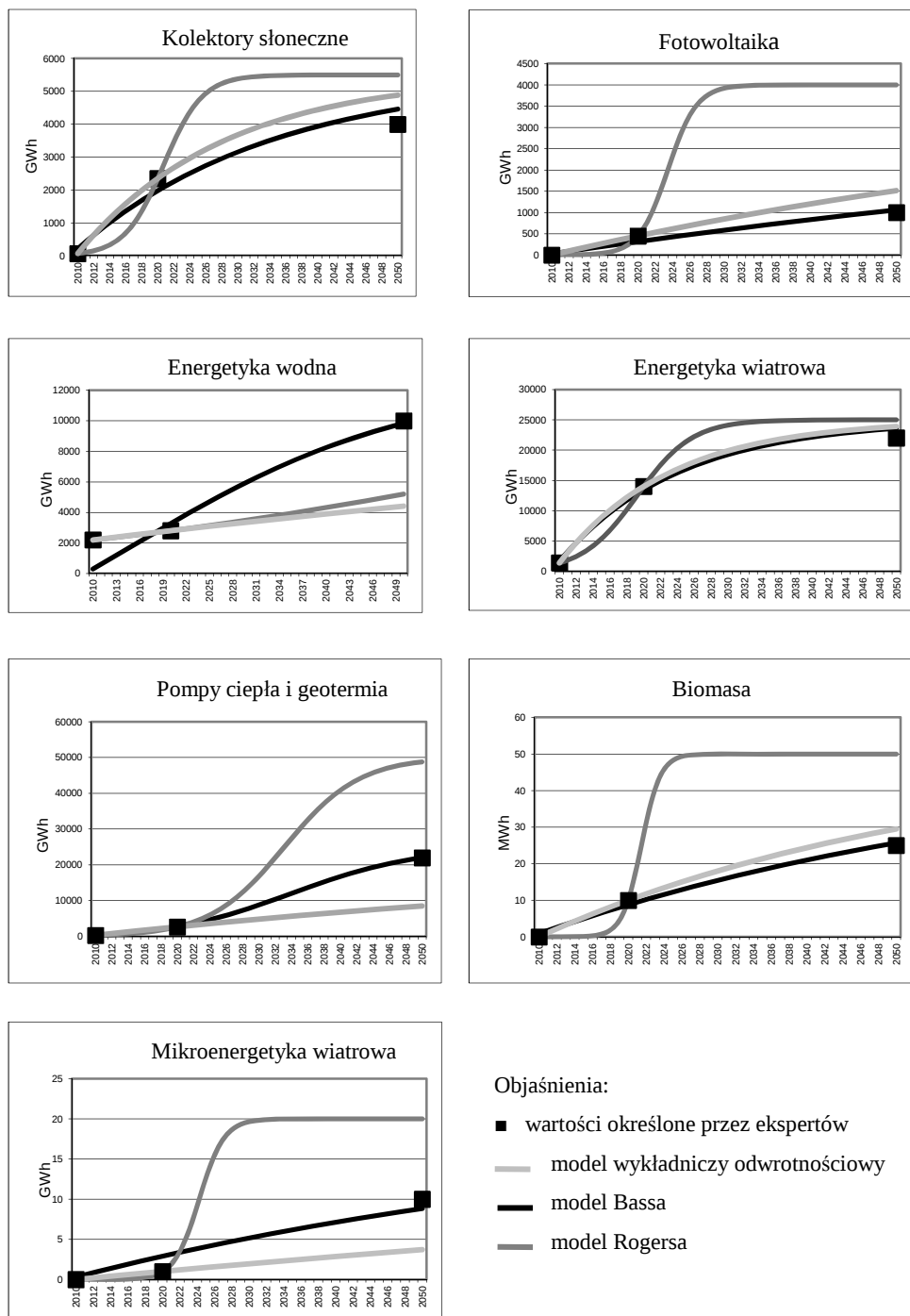
gdzie:

$y_{2010}, y_{2020}, y_{2050}$ – wartości określone przez ekspertów odpowiednio na lata 2010, 2020, 2050,

$\hat{y}_{2010}, \hat{y}_{2020}, \hat{y}_{2050}$ – wartości teoretyczne otrzymane na podstawie wstępnie oszacowanego modelu.

Wagi nadane poszczególnym odchyleniom przyjęto arbitralnie, chcąc w ten sposób nadać większe znaczenie ocenom ekspertów formułowanym na okresy bliższe teraźniejszości.

Prognozy rozwoju technologii OZE na lata 2010-2050, otrzymane na podstawie zbudowanych modeli przedstawiono na rys. 3.



Rys. 3. Prognozy rozwoju technologii OZE

Jako wada modelu Rogersa ujawniła się zbytnia wrażliwość na wartość potencjału rynkowego. Do oceny parametrów modelu, podobnie jak w przypadku modelu wykładniczego odwrotnościowego, nie wykorzystano wielkości produkcji energii określonej dla 2050 r. Odległość tej wartości od teoretycznego, wyznaczonego na podstawie modelu poziomu zjawiska może być pewnego rodzaju miernikiem dopasowania modeli do danych eksperckich. Na tej podstawie można stwierdzić, że najbliższy modelom jest (według ekspertów) wzrost produkcji energii wiatrowej. Powodem tego może być stosunkowo dobra znajomość tej technologii na tle innych, dopiero wkraczających na polski rynek. Zakładając, że intuicja ekspertów dotycząca przyszłości była trafna, można wnioskować, że technologie OZE nie będą się rozwijały zgodnie z logistycznym modelem dyfuzji innowacji.

W każdym z rozważanych przypadków model Bassa okazał się najlepiej dopasowany do posiadanych danych subiektywnych. Jedynie dla energetyki wodnej model ten, zwłaszcza w początkowym okresie, nie odpowiadał opiniom ekspertów. Można to tłumaczyć m.in. tym, że to źródło energii ma już za sobą fazę wdrożenia (wielkość produkcji energii wodnej w 2010 r. to 2200 GWh, a w ostatnich latach został zanotowany niewielki spadek tej wielkości), natomiast model dyfuzji Bassa służy do opisu rozwoju technologii dopiero pojawiających się na rynku, których aktualny poziom jest bliski zeru.

W badaniu *foresight* model przyczynowo-skutkowy Little'a (12) wykorzystano do wyznaczenia prognozy ilości zainstalowanej mocy w elektrowniach wiatrowych (Y). Jako główny czynnik wpływający na tę zmienną eksperci wskazali cenę energii elektrycznej (X). Bazowe (aktualne w czasie przeprowadzania badania) wielkości zmiennych występujących w modelu to: $Y = 1000$ MW, $X = 40$ zł. Ponadto, według opinii eksperta:

- w przypadku braku opłat za energię elektryczną, w okresie prognozy moc zainstalowana w elektrowniach wiatrowych utrzymywałaby się na poziomie 10 MW rocznie,
- przy nieograniczeniu dużych kosztach energii moc zainstalowana w elektrowniach wiatrowych mogłaby wynieść nawet 2000 MW,
- podwyższenie kosztów energii o 20% przyczyniłoby się do zwiększenia mocy w elektrowniach wiatrowych do 1300 MW.

Biorąc pod uwagę powyższe informacje oraz zależności dane wzorami (13), (14) otrzymano następujący model:

$$y_t = 10 + 1990 \frac{x_t^{3,42}}{297342 + x_t^{3,42}}$$

Ponieważ są dostępne dane historyczne dotyczące cen energii elektrycznej, wartości zmiennej objaśniającej w prognozowanych okresach mogą być określone na podstawie zbudowanego w tym celu prognostycznego modelu obiektywnego.

Podsumowanie

W obliczu wzrastającej popularności badań *foresight* zachodzi potrzeba konstrukcji długookresowych prognoz i scenariuszy rozwoju nowych technologii, o których brakuje obiektywnych danych z przeszłości. Podstawowym źródłem informacji są więc dane subiektywne pozyskane od ekspertów dziedzinowych. Na podstawie takich danych możliwa jest konstrukcja formalnego modelu prognostycznego – modelu subiektywnego. Postać funkcyjną modelu należy przyjąć *a priori*, wspomagając się przy tym np. sądami eksperckimi lub informacjami o technologiach analogicznych. Parametry modelu określa się na podstawie opinii ekspertów, zadając im w tym celu określone pytania.

Jeżeli znany jest przybliżony kształt zależności funkcyjnej zmiennej prognozowanej od zmiennej objaśniającej można tak sformułować pytania do ekspertów, aby otrzymać parametry modelu przyczynowo-skutkowego. Istotną zaletą takiego modelu jest możliwość uzupełnienia informacji subiektywnej o dane historyczne, które mogą być użyte do prognozowania wartości zmiennej objaśniającej.

Formalne modele subiektywne mogą być cennym narzędziem do konstrukcji długookresowych prognoz rozwoju nowych zjawisk w przypadku niepełnej informacji, wynikającej jedynie z częściowych opinii ekspertów o możliwym kształtowaniu się zjawiska. Badania wykazują [12], że takie modele mogą wygenerować bardziej trafne prognozy niż otrzymane bezpośrednio od ekspertów. Koniecznym warunkiem efektywności prognozowania jest w takiej sytuacji „dobra jakość” danych eksperckich, wymagających dogłębnej wiedzy eksperta na temat struktury prognozowanego zjawiska i rzetelności w procesie ich pozyskiwania.

Literatura

1. Armstrong J.S., Collopy F., *Integration of Statistical Methods and Judgment for Time Series Forecasting*, w: G. Wright, P. Goodwin, *Forecasting with Judgment*, John Wiley & Sons, New York 1998.
2. Bass F.M. (1969): *A new product growth model for consumer durables*, „Management Science” 1969, January.
3. Cohen W.A., *The Practise of Marketing Management. Analysis, Planning and Implementation*, Macmillan Publishing Company, New York 1991.
4. Dittmann P., *Prognozowanie w przedsiębiorstwie. Metody i ich zastosowanie*, Oficyna Ekonomiczna, Kraków 2004.
5. Gardner E., Jr., *Forecast with No Data*, „Lotus” 1991, Vol. 7, No. 6.
6. Lilien G.L., Rangaswamy A., *Marketing Engineering: Computer-Assisted Marketing Analysis and Planning, Revised Second Edition*, Trafford Publishing 2004.
7. Little J.D., *Models and Managers: The Concept of a Decision Calculus*, „Management Science” 1970, Vol. 16, No. 8.
8. Makridakis S., Wheelwright S.C., Hyndman R.J., *Forecasting: methods and applications*, J. Wiley, New York 1998.
9. Poradowska K., *Subjective Growth Models in Long-Term Forecasting the Development Technologies*, Econometrics – Forecasting, Research Papers, Uniwersytet Ekonomiczny, Wrocław 2011.
10. Poradowska K., *Wybrane aspekty prognozowania sprzedaży nowych produktów*, red. P. Miłobędzki, M. Szreder, Uniwersytet Gdański, Sopot 2011.
11. Rogers E.M., *Diffusion of Innovations*, Free Press, New York 1983.
12. Russo E.J., Schoemaker P.J., *Decision Traps: The Ten Barriers to Brilliant Decision-Making and How to Overcome Them*, Knopf Doubleday Publishing Group 1989.
13. Saunders J., *The specification of aggregate market models*, „European Journal of Marketing” 1987, Vol. 21, No. 2.
14. Shim J.K., *Strategic Bussines Forecasting*, St. Lucie Press, New York 2000.
15. Sokele M., *Growth models for the forecasting of new product market adoption*, „Elektronik” 2008, No. 3/4.

SUBJECTIVE MODELS IN THE DESIGN OF LONG-TERM FORECASTS

Summary

In the paper are presented some aspects of construction and practical use of subjective forecasting models. The main objective was to identify the usefulness of these models in the design of long-term forecasts and scenarios for the development of new

technologies. Theoretical considerations are supplemented with real examples of data analysis, obtained in the study of foresight "zero carbon energy economy in a sustainable development of the Polish to 2050", conducted by the Central Mining Institute in Katowice.

Włodzimierz Szkutnik

Uniwersytet Ekonomiczny w Katowicach

STATYSTYCZNA NIEOKREŚLONOŚĆ W WYCENIE CHARAKTERYSTYK RYNKÓW FINANSOWYCH

Wprowadzenie

Tematem rozważań będą modele charakterystyczne dla finansowego rynku w sytuacjach statystycznie nieokreślonych, finansowe kalkulacje w takich warunkach i zarządzanie portfelem pochodnych finansowych instrumentów. Ponadto odchodząc od typowego założenia dotyczącego stochastycznej struktury cen akcji przy modelowaniu stóp zwrotu (ich logarytmów) rozpatrzony zostanie wariantowy przypadek, w którym nie czyni się takich założeń. Rozważany też będzie „prosty” jednoetapowy model (B, S) -rynku, składający się z bankowego konta (pozbawiony ryzyka papier wartościowy) $B = (B_n)$ i jednej akcji $S = (S_n)$ z $n = 0, 1$. Oceniona zostanie górna i dolna gwarantowana cena w tym jednokrotnym modelu rynku.

Rozpatrzona w opracowaniu egzemplifikacja modelu finansowego rynku w warunkach statystycznej nieokreśloności została ujęta na przykładzie zadania modelowania racjonalnego zachowania inwestora. Pokazano sposób stwierdzenia istnienia dopuszczalnego rozwiązania zadania poszukiwania strategii zarządzania portfelem papierów wartościowych, zabezpieczającej zadany poziom dochodowości i ryzyka. Omówiono algorytm skonstruowania optymalnego zarządzania portfelem pochodnych instrumentów finansowych w warunkach nieokreśloności w określeniu dochodowości tego lub innego finansowego instrumentu.

1. Model finansowego rynku dla sytuacji stochastycznie nieokreślonych

W klasycznym sformułowaniu zadania wyboru struktury portfela ryzykownych inwestycji zakłada się istnienie N ryzykownych papierów wartościowych (lokat) z zadanymi dochodami (zyskiem) r_i , $i = 1, \dots, N$, które są traktowane jako zmienne losowe. Przyjmujemy następujące oznaczenia:

x_i – wartość oczekiwana i -tego papieru wartościowego,

$V = \{\sigma_{ij}\}$ – macierz kowariancji między i -tym a j -tym papierem wartościowym.

Założymy także, że istnieje możliwość lokowania pozbawionego ryzyka i odpowiednią zyskowność oznaczmy przez x_0 .

Portfel określony jest formalnie przez wektor:

$$(y_0, y)^T = (y_0, y_1, \dots, y_n) \in R^{N+1}$$

którego każda współrzędna odpowiada części kapitału zainwestowanego w odpowiedni finansowy instrument.

$y_0 = 1 - (e, y)$, gdzie symbol $(\cdot, \cdot)$ oznacza iloczyn skalarny; o wartości akcji y_i nie dopuszcza się ujemności, co stwarza możliwość otrzymania kredytu na nieryzykowny procent przy niedopuszczalnej krótkiej sprzedaży, co znaczy, że nie jest to możliwe w obecności instrumentów *short selking* (krótka sprzedaż).

Dochodowość portfela określa równość:

$$\mu(\bar{x}, \bar{y}) = y_0 \cdot x_0 + (y, x)$$

a odpowiadające tej dochodowości ryzyko:

$$\sigma(y) = (y^T \cdot V \cdot y)^{\frac{1}{2}}$$

gdzie:

$$\bar{x} = (x_0 \cdot x)^T, \quad \bar{y} = (y_0, y)^T$$

Charakterystyki:

$$\mu = \mu(\bar{x}, \bar{y}) \text{ i } \sigma = \sigma(y)$$

efektywnych (niedominujących) portfeli określone są w relacji (macierz V z założenia jest dodatnio określona):

$$\mu = x_0 + g \cdot \sigma, \quad g = \sqrt{(x - x_0 \cdot e)^T V^{-1} (x - x_0 \cdot e)}, \quad \mu \geq x_0$$

W tej relacji wartość x_i i σ_{ij} są z założenia z góry zadane i stałe.

Struktura efektywnego portfela

Przy zadanej, odpowiadającej zadanej dochodowości μ , poziomowi ryzyka σ , struktura efektywnego portfela może być wyrażona następująco:

$$y = \frac{V^{-1}(x - x_0 \cdot e)}{(x - x_0)^T V^{-1}(x - x_0 e)} (\mu - r_0)$$

lub odpowiednio:

$$y = \frac{V^{-1}(x - x_0 \cdot e)}{\sqrt{(x - x_0 e)^T V^{-1}(x - x_0 e)}} \cdot \sigma$$

1.1. Wariant zadania zachowania efektywności portfela w warunkach zmieniających się oczekiwanych dochodowości aktywów

Przy zmieniającej się oczekiwanej dochodowości aktywów, co jest oczywiste w warunkach rynkowych, założymy, że wielkość x_i reprezentująca ryzykowne papiery wartościowe (lokaty) jest zmienna w czasie i jej dynamikę opisuje równanie różniczkowe:

$$\frac{dx}{dt} = A(t) \cdot x + q(t), \quad q(t) \in Q(t), \quad t_0 \leq T \leq \theta \quad (1)$$

$$x(t_0) = x^0 \quad (2)$$

Wieloznaczna funkcja $Q(t)$ (zbiór) wyraża nieokreśloność w zmianach dochodowości. Będziemy zakładać, że zobrazowanie funkcji $Q(t)$ jest kawałkami ciągle względem t i przy każdym ustalonym t zbiory $Q(t)$ są wypukłe i zwarte, przy czym $0 \in Q(t)$.

Zakłada się, że można zmieniać strukturę portfela w każdym momencie $t \in [t_0, \theta]$ z ograniczoną dynamiką (tempem, prędkością). Dynamika zmienności portfela w tym przypadku jest opisana w równaniu:

$$\frac{dy}{dt} = u, \quad y_0 = 1 - (e, y)$$

gdzie u oznacza działanie zarządcze, ograniczające wpływ efektywnościowy wyrażone w relacji $u \in P(t)$; zakłada się, że $P(t)$ ma własności analogiczne do własności zobrazowania $Q(t)$.

1.2. Cel inwestycyjnej korekty działań efektywnościowych – – zachowania efektywności

Oczywiste jest, żeby wymagać od oddziaływania efektywnościowego u zachowania efektywności portfela i zabezpieczenia jego charakterystyk: ryzyko-dochodowość. Merytoryczne sformułowanie zadania zakłada wybór strategii zarządczej realizującej się w postaci zarządczego wpływu efektywnościowego działania $u = u(t)$, które formułuje się według dostępnej w momencie t informacji i zabezpiecza wymagane własności portfela. Określona strategia wyraża regułę określającą dynamikę wektora Y rozkładu zainwestowanych środków w ryzykowne i wolne od ryzyka aktywa.

Formalnie dopuszczalną strategią regulującą skład portfela dla podtrzymywania jego efektywnościowych cech, określane jest wieloznaczne przekształcenie:

$$U(t, x, y)$$

mierzone względem t , półciągłe z góry względem x, y oraz z wypukłymi zwarzającymi wartościami:

$$U(t, x, y) \subseteq P(t).$$

1.3. Dynamiczna restrukturyzacja portfela

Formalna procedura w dynamicznej restrukturyzacji portfela ma następujący przebieg:

Równanie (1) i równanie:

$$\frac{dy}{dt} = u, \quad u \in U(t, x, y) \quad (3)$$

przy uczynionych powyżej założeniach ma bezwzględnie ciągłe rozwiązanie:

$$x(t), y(t), \quad \text{gdzie} \quad y_0(t) = 1 - (e - y(t)), \quad t_0 \leq t \leq \theta$$

dla dowolnych początkowych warunków (2) oraz:

$$y(t_0) = y_0, \quad y_0^0 = 1 - (e, y^0) \quad (4)$$

przy czym rozwiązanie przedłużone jest na cały odcinek czasu $[t_0, \theta]$.

Dla każdego rozwiązania wartości $x(t)$ i udziałów $y(t)$:

$$x(t) = (x_0(t), x(t))^T, \quad y(t) = (y_0(t), y(t))^T$$

można rozpatrzeć ewolucję charakterystyk portfela, wyrażających dochodowość i ryzyko:

$$\mu(t) = y_0(t) \cdot x_0(t) + (y(t), x(t))$$

$$\sigma(t) = (y^T(t) \cdot V^{-1} \cdot y(t))^{\frac{1}{2}}$$

Zbiorowi efektywnych portfeli na płaszczyźnie μ, σ odpowiada prosta zmieniająca się w czasie. Dla efektywności portfela wystarcza, aby $\mu(t)$ i $\sigma(t)$ były ze sobą w analogicznej relacji, jak wcześniej charakterystyki efektywnych (niedominujących) portfeli, tzn.

$$\mu(t) = x_0(t) + g(t) \cdot \sigma(t)$$

gdzie:

$$g(t) = \sqrt{(x(t) - x_0(t) \cdot e)^T V^{-1} (x(t) - x_0(t) \cdot e)} = \|x(t) - x_0(t) \cdot e\| \cdot V^{-1}$$

2. Finansowe oceny (kalkulacje) na rynku dla sytuacji nieokreślonych statystycznie

Stawiamy następujący cel regulacyjnych działań efektywnościowych – zachowanie efektywności portfela w rozpatrywanym przedziale czasu:

Założmy, że stopa wolna od ryzyka jest stała:

$$x_0(t) = r_0, \quad \text{tj.} \quad \frac{dx_0}{dt} = 0, \quad \text{oprócz tego} \quad A(t) \equiv 0$$

Wtedy dochodowość portfela wynosi:

$$\mu(x(t), y(t)) = y_0(t) \cdot r_0 + (x(t), y(t))$$

a odpowiednie ryzyko:

$$\sigma(y(t)) = \sqrt{(y^T(t) \cdot V \cdot y(t))} = y(t)_V$$

Dynamiczną efektywność portfela zabezpiecza się według formuły wyrażonej we wzorze:

$$\mu(x(t), y(t)) = r_0 + g(t) \cdot \sigma(T)$$

Można także wskazać postać działania efektywnościowego wyrażonego przez $u(t)$, istnienie dopuszczalnej strategii-działania efektywnościowego przy zarządzaniu portfelem papierów wartościowych, która zabezpiecza zadany poziom dochodowości i ryzyka. Ponadto można wskazać model zachowania inwestora na rynku instrumentów finansowych oraz optymalną strategię zarządzania i kryterium optymalności.

Przyjmujemy, że y^0 – efektywny portfel dla $\hat{x}^0$ z oczekiwaną wartością dochodowości:

$$\mu^0 = \mu^0(x^0, y^0) \text{ i } \sigma^0 = \sigma^0(y^0)$$

2.1. Dopuszczalna strategia zarządzania

Problem, jaki się pojawia w zaprezentowanym wyżej ujęciu efektywności portfela wiąże się ze wskazaniem strategii zarządzania-działania efektywnościowego $U = U(t, x, y)$, będącej „gwarantem” efektywności portfela $y(t)$ dla $x(t)$ ($t_0 \leq t \leq \theta$), jakie by nie były rozwiązania $x(t)$, $y(t)$ równań różniczkowych (1), (2) i (3), (4).

Dopuszczalna strategia zachowania efektywności portfela

W pierwszej kolejności skupimy uwagę na wskazaniu dopuszczalnego działania efektywnościowego, będącego rozwiązaniem powyższego problemu i zabezpieczającego nałożony poziom ryzyka:

$$\sigma(y(t)) \leq \sigma^0$$

Ponadto wskazana zostanie strategia-działanie efektywnościowe, będące rozwiązaniem sformułowanego problemu i zabezpieczającego wymagany poziom dochodowości:

$$\mu(x(t), y(t)) \geq \mu^0$$

W dalszej części będziemy zakładać spełnienie następującego warunku regularności:

$$\|x^0 - e \cdot r_{-0}\|_{V^{-1}} - \max_{\|l\|_{V^{-1}}=1} \int_{t_0}^{\theta} \rho(l/Q(t))dt = d > 0 \quad (5)$$

gdzie:

$\rho(l/Q(t))$ – funkcja oporowa zbioru $Q(t)$,

$\rho(l/z) = \max\{(l, z) : z \in Z\}$.

W kontekście omawianego tu istnienia działania efektywnościowego ważne jest to, że spełnienie warunku regularności (5) jest istotne dla rozwiązania sformułowanego problemu, a ponadto uwzględnia występowanie na rynku pożądanych sytuacji. Mamy bowiem własność:

1. Warunek regularności zabezpiecza relacje:

$$\|x(t) - r_0 \cdot e\|_{V^{-1}} \geq d \quad (6)$$

dla każdego $t \in [t_0, \theta]$, co oznacza, że niemożliwa jest sytuacja, kiedy dochodowości wszystkich ryzykownych aktywów są jednocześnie bliskie stopie wolnej od ryzyka.

Główny problem rozpatrywany w tym rozdziale dotyczący wskazania dopuszczalnej strategii-działania efektywnościowego portfela i jednocześnie zapewniającej narzucony poziom dochodowości μ^0 jest rozwiązany poprzez wskazanie relacji zachodzącej między wieloznacznymi przekształceniami $Q(t)$ i $P(t)$. Warunek ten, który jest wystarczający, jest wyrażony następująco:

$$d^2 \cdot \rho(l/P(t)) - (\mu^0 - r_0) \cdot \rho(l/V^{-1}/Q(t)) \geq (\mu^0 - r_0) \max_{q=1} \rho(q/Q(t)) \quad (7)$$

$$\forall_l l_{V^{-1}} = 1, \forall_t t \in [t_0, \theta]$$

Natomiast wskazanie dopuszczalnej strategii-działania efektywnościowego zapewniającego nieprzekroczenie poziomu ryzyka σ^0 wymaga spełnienia relacji:

$$d \cdot \rho(l/P(t)) - \sigma^0 \cdot \rho(l/v^{-1}Q(t)) \geq \sigma^0 \cdot \max_{q=1} \rho(q/Q(t)) \quad (8)$$

W dalszej kolejności rozważania prowadzą do określenia zbioru wszystkich możliwych wartości określających dynamikę struktury efektywnego portfela.

2.2. Krańcowa struktura efektywności portfela jako element działań efektywnościowych

Okazuje się, czego nie będziemy tu wykazywać, że przy spełnieniu warunku regularności (5) i dostatecznego warunku regularności (8) zapewniającego zadany poziom dochodowości μ^0 , a także wtedy, gdy spełnione są warunki (5) oraz (7) zachodzi inkluzja:

$$R(t) \subseteq P(t)$$

dla prawie wszystkich $t \in [t_0, \theta]$.

Zbiór $R(t)$ dla każdego t wyraża możliwe wartości krańcowej struktury efektywności portfela, co jest tożsame z wyrażeniem dynamiki struktury efektywnego portfela:

$$\frac{dz(t)}{dt}$$

Wartości te odpowiadają możliwym realizacjom $z(t)$ określającym strukturę efektywnego portfela, odpowiadającego zadanej dochodowości lub zadanemu poziomowi ryzyka σ . Realizacje $z(t)$ określone są bowiem następująco:

$$z(t) = \frac{V^{-1}(x(t) - r_0 \cdot e)}{\|x(t) - r_0 \cdot e\|_{V^{-1}}}$$

Jeśli wprowadzimy oznaczenia $\tilde{x}(t) = x(t) - r_0$, to $z(t) = \frac{V^{-1} \cdot \tilde{x}(t)}{\|\tilde{x}(t)\|_{V^{-1}}} \cdot \sigma^0$.

Istnienie dopuszczalnej strategii

Jeżeli spełniony jest warunek regularności (5) i warunek dostateczny (8), to istnieje dopuszczalna strategia-dopuszczalne działanie efektywnościowe będące rozwiązaniem problemu istnienia dopuszczalnej strategii gwarantującego efektywność portfela $y(t)$ dla $x(t)$ $t_0 \leq t \leq \theta$ przy dowolnych rozwiązaniach $x(t)$, $y(t)$ równań (1), (2) i (3), (4).

Prześledzenie istnienia takiej sytuacji jest bardzo przydatne dla wnikięcia w istotę omawianego ujęcia zagadnienia korygowania portfela akcji dla zapewnienia jego wymaganej efektywności.

Określimy wektor S_0 , w postaci znormalizowanej:

$$S_0 = \begin{cases} \frac{z - y}{\|z - y\|}, & z \neq y \\ 0 & z = y \end{cases}$$

Wtedy np. dla $(z - y) = (0,1; 0,9)$, $\|z - y\| = \sqrt{0,01 + 0,81} = \sqrt{0,81}$

$$\text{ i } S_0 = \left(\frac{0,1}{\sqrt{0,82}}; \frac{0,9}{\sqrt{0,82}} \right) \text{ i } \|S_0\| = \sqrt{\frac{0,01}{82} + \frac{0,81}{82}} = 1.$$

Wieloznaczną funkcję można określić w postaci:

$$u^e(t, x, y) = \{u^e \in U(t) : (S_0, U^e) = \max_{u \in P}(S_0, u)\}$$

Funkcja ta jest półciągła z góry i zależy od z i y , które z kolei jako ciągłe (w sposób ciągły) zależą od x i y w obszarze:

$$\|\tilde{x}\|_{V^{-1}}^2 > 0$$

gdz:

$$z = \frac{V^{-1} \cdot \tilde{x}}{\|\tilde{x}\|_{V^{-1}}} \text{ lub } z = \frac{V^{-1} \tilde{x}}{\|\tilde{x}\|_{V^{-1}}^2} (\mu^0 - r_0)$$

gdzie:

$\tilde{x} = x - r_0 \cdot e$ – wektor wartości akcji w portfelu bez papierów wartościowych wolnych od ryzyka.

Z określenia U^e i własności przekształcenia $P(t)$ wynika mierzalność względem t strategii i tym sposobem określona wyżej strategia jest dopuszczalna. Należy tu w sposób zasadniczy zaakcentować, że z formalnego punktu widzenia dopuszczalną strategią jest wieloznaczne przekształcenie $\mu = u(t, x, y)$ mierzalne względem t itd. Można w tym kontekście wykazać, że strategia ta realizuje postulat dochodowości μ^0 . W tym celu przyjmuje się, że $x(t)$, $y(t)$ – rozwiązania równań (1), (2) i (3), (4), gdzie w warunku (3) $U(t, x, y) = U^e(t, x, y)$.

Dla prawie wszystkich $t \in [t_0, \theta]$ mamy oszacowanie:

$$\begin{aligned} \frac{d}{dt} ((z(t) - y(t))^2) &= 2 \cdot ((z(t) - y(t), (z'(t) - u^e(t))) = \\ &= 2[(z(t) - y(t), z'(t)) - \max_{u \in P(t)} (z(t) - y(t), u)] \leq \\ &\leq 2[\max_{z' \in R(t)} (z(t) - y(t), z'(t)) - \max_{u \in P(t)} (z(t) - y(t), u(t))] = \\ &= 2[\rho(l(t) / R(t)) - \rho(l(t) / P(t))], \end{aligned}$$

gdzie:

$$l(t) = z(t) - y(t).$$

Wystarczy jeszcze zauważyć, że jeśli:

$$z = \frac{V^{-1} \tilde{x}(t)}{\|\tilde{x}(t)\|_{V^{-1}}} \sigma^0$$

to z (6) i (8) otrzymujemy $\frac{d}{dt} ((z(t) - y(t))^2 \leq 0$. Wynika to także stąd, że $R(t) \subseteq P(t)$. Stąd i z równości $z(t_0) = y(t_0)$ wynika, że:

$$y(t) \equiv z(t)$$

Portfel jest zatem efektywny i zabezpiecza zadany poziom ryzyka σ^0 .
Jeśli:

$$z = \frac{V^{-1}(\tilde{x}(t))}{\|\tilde{x}(t)\|_{V^{-1}}^2} (\mu^0 - r_0)$$

to z równań (6) i (7) znowu otrzymujemy:

$$\frac{d}{dt} (z(t) - y(t))^2 \leq 0$$

Portfel jest zatem efektywny i zabezpiecza zadany poziom dochodowości μ^0 . Dzieje się to w czasie $t_0 \leq t \leq \theta$; portfel $y(t)$ dla $x(t)$ jest efektywny, co gwarantuje dopuszczalna strategia $U = U(t, x, y)$.

3. Dynamika cen aktywów w systemie zabezpieczeń w warunkach nieokreśloności

Stosując terminologię teorii systemów sformułujemy problem oceny fazowego wektora liniowego systemu, funkcjonującego w obecności wymuszeń (sterowań) i błędów (pomiarów). Terminologia ta i sam problem zostaną zaadaptowane na oceny dynamiki cen aktywów.

3.1. Specyfika zadań oceny

Powszechnie wiadomo, zarówno z literatury fizycznej (technicznej), jak i dotyczącej ocen dokonywanych dla wielkości rynku finansowego, że specyfika zadań oceniania ściśle łączy się z charakterem informacyjnych założeń, przy których są one rozwiązywane. Różnorodność problemów i ich sformalizowania prowadzi do wielu ujęć w ich rozwiązywaniu. W tej części rozdziału głównym założeniem będzie przyjęcie, że w ocenie dynamiki cen aktywów rozpatrywanej w układzie zabezpieczeń możliwe jest określenie problemu i jego rozwiązanie z różnych punktów widzenia odnośnie do warunków nieokreśloności.

Istnieją różne aspekty tych warunków:

1. Warunki stochastyczne określone.

Jeśli istnieje pełny statystyczny opis nieznanymi zakłóceń i początkowych danych, to rozwiązanie osiąga się w ramach stochastycznej teorii obserwacji lub filtracji.

2. Statystyczne warunki nieokreśloności.

Z warunkami nieokreśloności łączą się zadania, w których statystyczny opis wskazanych apriorycznych danych w ogóle nie występuje. Wiadomości o nich ograniczają się jedynie do zadania dopuszczalnych obszarów zmian nieznanymi wielkościami. W tym przypadku na podstawie teorii obserwacji w warunkach nieokreśloności możliwe jest osiągnięcie rozwiązania z wykorzystaniem teorii gier.

3. Warunki mieszane

Z takim zespołem warunków spotykamy się najczęściej np. przy tytułowym problemie oceny dynamiki cen aktywów. W takim przypadku informacja apriori o wymuszeniach i błędach ma mieszany charakter. A właściwie, w takich przypadkach rozpatrywane są zadania oceniania, kiedy w systemie i także w kanale pomiaru współ ze stochastycznymi istnieją wymuszenia, błędy, o których w ogóle brakuje jakiegokolwiek statystycznej informacji, a informacja o nich ogranicza się tylko do opisu pewnych dopuszczalnych obszarów ich zmian. Tak jest w układzie warunków np. w wypadku cen akcji. Ponieważ zmieniają się one z reguły skokowo, a nie ciągle, uzasadnione jest ustalenie ich dopuszczalnego

poziomu zmian. Takie sytuacje, gdy jednocześnie występują wymuszenia jako stochastyczne ze znanymi charakterystykami statystycznymi, jak i wymuszenia, informacje o których ogranicza się do zadania obszaru ich zmian nazywa się statystycznie nieokreślonymi.

W dalszej części przedmiotem rozważań będą zagadnienia, w których rozpatrywane będą zadania oceny w warunkach nieokreśloności i statystycznie nieokreślonych. Podstawą przyjętego w rozdziale podejścia do statystycznej nieokreśloności zadań oceniania były teoriogrowe metody rozwiązywania zadań sterowania i obserwacji, rozwinięte przez N.N. Krasowskiego [3]. Konkretnie sformułowania zadań są związane z koncepcją aposteriori obserwacji. Niektóre z nich przytoczone zostały w modelowych przykładach.

3.2. Model dynamiki zmienności cen w warunkach statystycznie nieokreślonych

Niech w liniowym przybliżeniu dynamika zmiany cen opisana jest różnicowym układem równań:

$$x_{K+1} = A_K \cdot x_K + B_K \cdot W_K + C_K \cdot \xi_K, \quad k = 0, 1, \dots \quad (9)$$

Zakłada się, że w momencie czasu $k = 1, \dots$ jest znana informacja w postaci wektora $y_K \in R^n$ odpowiadającego wektorowi fazowych współrzędnych $x_K \in R^n$ w liniowej równości:

$$y_K = G_K \cdot x_K + H_K \cdot V_K + \eta_K, \quad k = 0, 1, \dots \quad (10)$$

Tutaj w (9) i (10) ξ_K, η_K są odpowiednio wymuszeniami i błędami obserwacji (błędy informacyjne) – niezależne gaussowskie ciągi, przy czym przyjmujemy tu, że:

$$E\xi_K = 0, \quad E\eta_K = 0, \quad \text{cov}\{\xi_K\} = Q_K, \quad \text{cov}\{\eta_K\} = R_K$$

gdzie $Q_K \in R^{mm}$, $R_K \in R^{rr}$ są dodatnio określonymi macierzami kowariancji.

O deterministycznych oddziaływaniach $w_K \in W_K \subset R^m$ i błędach obserwacji $v_K \in V_K \subset R^p$ wcześniej nie czyni się żadnych założeń. Natomiast macierze A_K, B_K, C_K, G_K, H_K o odpowiednich wymiarach i wypukłe zbiory W_K, V_K traktuje się jako znane.

Początkowy stan:

$$x_o = x_{co} + x_{no}$$

układu (9) jest traktowany jako n -wymiarowy wektor gaussowski, niezależny od ξ_K, η_K , o znanej dodatnio określonej macierzy kowariancji $\text{cov}(x_{co}) = P_o$, ale z nieznaną wcześniej średnią wartością $x_{no} = Ex_o = \tilde{x}_o \in \bar{X}_o$ – znany wypukły zbiór zwarty.

Podsumowując, w równaniach dynamiki zmian cen i obserwacji występują jako losowe wektory:

$$x_{co}, \xi_K, \eta_K$$

których charakterystyki statystyczne są w pełni znane, jak i nieokreślone wektory:

$$x_{no}, V_K, W_K,$$

o których informacja ogranicza się tylko do ich przynależności do obszarów zmian, odpowiednio:

$$x_{no} \in \bar{X}, w_K \in W_K, v_K \in V_K \quad (11)$$

Ograniczenia typu (11) nazywane są chwilowymi lub nagłymi, a jeszcze inaczej geometrycznymi ograniczeniami.

Zadania oceny

W dalszej części będą rozpatrywane tylko zadania oceny, dlatego w układzie (9) nie występuje sterowanie, więc $u_K = 0$. Dla każdej ustalonej pary ciągów:

$$w_{N-1}(\cdot) = \{w_o, \dots, w_{N-1}\}, v_N(\cdot) = \{v_1, \dots, v_N\}$$

wielkości:

$$\bar{\xi}_K = B_K \cdot w_K, \bar{\eta}_{K+1} = H_K \cdot v_{K+1}$$

są w istocie średnimi wartościami gaussowskich ciągów:

$$\xi_K^* = B_K \cdot w_k + C_K \cdot \xi_K, \mu_{K+1}^* = H_{K+1} v_{K+1} + \eta_{K+1}, \quad k = 0, \dots, N-1$$

Statystyczna nieokreśloność modelu

Z powyższego wynika, że równania (9), (10) mogą być rozpatrywane jako równania z gaussowskimi wejściowymi oddziaływaniami ξ_K^* , η_K^* i gaussowskim wektorem początkowym x_0 . Statystyczna nieokreśloność wynika z tego, że o średnich wartościach wiadomo tylko tyle, ile wynika z (11). Jest tak dlatego, że zbiór apriorycznych rozkładów z niezależnych wzajemnie gaussowskich rozkładów o znanych wcześniej macierzach kowariancji P_0, Q_k, R_{k+1} , $k = 0, \dots, N-1$, ale nieznanymi średnimi wartościami, informacja o których ogranicza się poprzez zadane warunki (11).

Aprioryczne rozkłady są w pełni określone poprzez wybór wielkości:

$$\xi_N(\cdot) = \{x_0, w_{N-1}(\cdot), v_N(\cdot)\}$$

rozpatrywanej jako element przestrzeni $R^n \times R^{mN} \times R^{PN}$.

Można tu zauważyć, że jeśli zaobserwowany byłby ciąg obserwacji:

$$y_N(\cdot) = \{y_1, \dots, y_N\}$$

oraz dana byłaby wielkość $\xi_N(\cdot)$, to najlepszą oceną dla parametrów rozkładu wektora byłaby warunkowa wartość oczekiwana:

$$E[x_N / y_{N(\cdot)}] = x_N^*$$

Ponadto dla aposteriori macierzy kowariancji:

$$P_N = E[(x_N - x_N^*)(x_N - x_N^*)' / y_{N(\cdot)}]$$

można sformułować stosowne równanie Riccatiego.

Dokonując przeglądu wszystkich możliwych wielkości $\xi_N(\cdot)$ można otrzymać zbiór ocen x_N^* , które razem z macierzą P_N stanowią najlepszą informację o wektorze x_N , którą tylko można wyciągnąć na podstawie opracowania bieżącego poziomu cen $y_{N(\cdot)}$ na dany konkretny moment czasu.

Uzyskanie punktowej oceny oznacza teraz wybór punktu ze wskazanego zbioru, będącej najlepszą w sensie pewnego kryterium.

Wybór punktowej oceny metodą minimaksowo-stochastycznej filtracji

Rozpatrzmy w przestrzeni $R^n \times R^{mN} \times R^{PN}$ następujący zbiór:

$$D_K = \{\xi_K(\cdot) : \bar{x}_o \in \bar{x}_o, w_K \in W_K, v_{K+1} \in V_{K+1}, k = 0, \dots, N-1\}$$

Wtedy przyjmując jako funkcję strat:

$$r(x) = x^T x = \|x\|^2$$

dochodzimy do następującego zadania minimaksowo-stochastycznej filtracji:

Względem znanej realizacji obserwacji $y_K(\cdot)$ ocenę $x_N^* = x_N^*(y_N(\cdot))$ można wyznaczyć z warunku:

$$\begin{aligned} & \max_{\xi_N(\cdot) \in D_N} E\{\|x_N - x_N^*\|^2 / y_N(\cdot), \xi_N(\cdot)\} = \\ & = \min_{C \in R^n} \max_{\xi_N(\cdot) \in D_N} E\{\|x_N - c\|^2 / y_N(\cdot), \xi_N(\cdot)\} = \varepsilon_N^2 \end{aligned}$$

Symbol $E\{\cdot / y_N(\cdot), \xi_N(\cdot)\}$ oznacza warunkową wartość oczekiwaną przy ustalonej wartości $\xi_N(\cdot)$.

3.3. Minimaksowo-stochastyczna ocena opcji w statystycznym układzie

Śledząc minimaksowo-stochastyczne podejście reprezentowane w poprzednim podrozdziale przy ocenie opcji przyjmiemy obecnie, że x jest gausowską skalarną wielkością, tj. $x \sim N(\bar{x}, \sigma_x^2)$, przy czym σ_x^2 – średni kwadrat odchylenia (wariancji) zysku jest znany, a $\bar{x}$ nieznany. Wiadomo tylko, że zadany jest przedział:

$$|\bar{x}| \leq \alpha$$

gdzie α – znana stała.

Obserwowana jest skalarna wielkość:

$$z = n + \eta + v$$

gdzie $\eta \sim N(0, \sigma_\eta^2)$, przy czym $\sigma_\eta > 0$ jest znaną wielkością nazywaną wskaźnikiem niestabilności ceny akcji, a $|v| \leq \beta$ (β – znana stała).

Przypadek stałego $\xi(\cdot)$

Ustalając pewną wartość $\xi(\cdot) = (\bar{x}, v)$ można stwierdzić, że aposteriori rozkład wielkości x przy uwzględnieniu realizowanego wyniku pomiaru z jest gausowskim rozkładem:

$$x \sim N(\bar{x}, \sigma_{x/z}^2)$$

przy czym:

$$\begin{aligned} y &= \bar{x} + \sigma_{x/z}^2 \cdot \sigma_{\eta}^{-2} \cdot (z - \bar{x} - \bar{y}) \\ \sigma_{x/z}^2 &= \sigma_x^{-2} + \sigma_{\eta}^{-2} \end{aligned} \quad (12)$$

Jeśli przyjmiemy, że $\sigma_x = \sigma_y = 1$, wtedy $\sigma_{x/z}^2 = 0,5$

i $x^* = 0,5 \cdot \bar{x} + 0,5(z - v)$. Będzie to jednocześnie najlepsza ocena przy znanym wektorze $\xi(\cdot) = (\bar{x}, v)$.

Przy takich ustaleniach, jak wyżej przyjęto, że straty (średniokwadratowy błąd) wynoszą:

$$E = \{(x - y)^2 / z, \xi(\cdot)\} = 0,5$$

Przypadek nieznanego $\xi(\cdot)$

W przypadku, gdy $\xi(\cdot)$ jest nieznanymi, tzn. nieznaną jest średnia cena $\bar{x}$ i błąd obserwacji v , to przyjmując jako ocenę liczbę C ryzykujemy, że straty mogą osiągnąć wielkość:

$$\begin{aligned} \varphi(c) &= \max_{\xi(\cdot)} E\{(x - c)^2 / z, \xi(\cdot)\} = \\ &= \max_{\substack{|\bar{x}| \leq \alpha \\ |v| \leq \beta}} E\{(x - y + y - c)^2 / z, \xi(\cdot)\} = \max_{\substack{|\bar{x}| \leq \alpha \\ |v| \leq \beta}} [\sigma_{x/z}^2 + (x - c)^2] \end{aligned}$$

Podstawiając tutaj wyrażenie z (4) i wyznaczając maksimum otrzymujemy:

$$\varphi(c) = \begin{cases} \frac{1}{2} + \left(\frac{\alpha + \beta}{2} + \frac{1}{2}z - c\right)^2 & \text{przy } c \leq \frac{1}{2}z \\ \frac{1}{2} + \left(-\frac{\alpha + \beta}{2} + \frac{1}{2}z - c\right)^2 & \text{przy } c > \frac{1}{2}z \end{cases}$$

Zatem minimum z (4) przyjmie postać:

$$\min_{c \in R^n} \varphi(c) = \max_{\substack{|\bar{x}| \leq \alpha \\ |v| \leq \beta}} E\{(x - c)^2 / z, \bar{x}, v\} = \frac{1}{2} + \left(\frac{\alpha + \beta}{2}\right)^2$$

i to minimum osiąga się przy $c = x_0 = 0,5z$.

Podsumowanie

Rozpatrzone wybrane modele finansowego rynku w warunkach statystycznej nieokreśloności na przykładzie modelowania racjonalnego zachowania inwestora wykazały, że istnieje dopuszczalne rozwiązanie zadania o znajdowaniu strategii efektywnego zarządzania portfelem akcji i innych instrumentów finansowych, zabezpieczające zadany poziom dochodowości i ryzyka. W innych badaniach można spróbować stwierdzić istnienie algorytmicznego sposobu strategicznego zarządzania portfelem pochodnych instrumentów finansowych w warunkach nieokreśloności przy określonej dochodowości wybranego instrumentu finansowego. Należy tu zauważyć, że zasygnalizowane tematy są rzadko podejmowane w pracach polskich badaczy i warto je szerzej propagować.

Literatura

1. Dothan M.V., *Prices in Financial Markets*, Oxford Univ. Press, Oxford 1990.
2. Gołębiowski D.J., Dołmatow A.S., *Sterowanie portfelem pochodnych instrumentów finansowych*. RAN. „Teoria i Systemy Sterowania”, nr 4, Moskwa 2000.
3. Krasowski N.N., *Sterowanie i stabilizacja przy deficycie informacji*, RAN, „Techniczna Cybernetyka”, No. 1, Moskwa 1993.
4. *Podstawy stochastycznej finansowej matematyki*, FAZIS, T. 1, *Fakty. Modele*, T. 2, Moskwa 1998.

THE STATISTICAL UNCERTAINTY IN THE MEASUREMENT OF THE CHARACTERISTICS OF FINANCIAL MARKETS

Summary

Considered in developing the financial model exemplification of the market in terms of statistical uncertainty was recognized as an example the task of modeling rational investor behavior. Shows how to determine the existence of an acceptable solution of the problem search strategy portfolio management, hedging given level of profitability and risk. Discusses the algorithm to construct the optimal portfolio management, derivative financial instruments under conditions of uncertainty in determining the profitability of this or any other financial instrument.

Jerzy Zemke

Uniwersytet Gdański

RYZIKO W ASPEKCIE ZARZĄDZANIA W ZRÓŻNICOWANYM OTOCZENIU SPOŁECZNO-GOSPODARCZYM

Wprowadzenie

Praktyka zarządzania organizacją gospodarczą pomimo dynamicznego rozwoju staje przed zagadnieniem turbulencji uwarunkowań otoczenia, które wpływają na jakość procesów związanych z realizacją celów strategicznych. Otoczenie organizacji gospodarczej nie jest jednorodne. Teoria zarządzania wyróżnia mikro oraz makrootoczenie^{*}. Podział nie jest przypadkowy, powodowany jest właściwościami uwarunkowań otoczenia. Otoczenie makro tworzą uwarunkowania, które są identyfikowane z funkcjami: politycznymi, społecznymi, kulturowymi, demograficznymi, prawnymi, ekonomicznymi oraz technologicznymi, natomiast zakres mikrootoczenia wyznaczają szeroko pojmowane zasoby organizacji. Otoczenie, w którym funkcjonuje organizacja gospodarcza nie jest stabilne. Globalizacja gospodarki ułatwia przenoszenie zmian na rynki lokalne, co powoduje destabilizację uwarunkowań otoczenia, przy czym rezultat zmian może sprzyjać realizacji planów gospodarczych bądź stanowić istotne utrudnienie w procesach osiągania celów przyjętych w planach rozwoju organizacji gospodarczych. Zmiany uwarunkowań mają charakter losowy, a niepewność co do jakości i terminowości realizowanych celów wiąże się z będącym w powszechnym użyciu – nie tylko naukowym – pojęciem „ryzyka”. Pojęcie to funkcjonuje w literaturze naukowej od ponad 100 lat i doczekało się wielu definicji^{**}. Spośród wielu zamieszczanych w literaturze problemu cenię szczególnie tę, którą sformułował K. Jajuga [10], łącząc ryzyko z procesem decyzyjnym, a dokładniej z działaniami związanymi z realizacją tego procesu. Tak pojmowane ryzyko ma wyrazisty kontekst praktyczny wymagany w realizacji procesów zarządzania, ponieważ umożliwia rozwiązanie zagadnienia identyfikacji ryzyka, poczynając od rozwiązania problemu jego pomiaru, a kończąc na konstrukcji instrumentów wspomagania procesów zarządzania.

^{*} Podział oraz oddziaływanie otoczenia organizacji zob. [1, s. 102 i nast.].

^{**} Po raz pierwszy zdefiniował je A.H. Willet: [9, s. 6].

Obszerny kontekst ujęcia ryzyka wykracza poza objętość tego opracowania. Z tego powodu jego zawartość została ograniczona do przedstawienia źródeł ryzyka, zagadnienia identyfikacji ryzyka realizacji procesów decyzyjnych oraz pomiaru ryzyka. Szczegółne miejsce zajmuje zagadnienie źródeł ryzyka oraz jego identyfikacja. Teoria zarządzania organizacją gospodarczą zwraca uwagę na tzw. zmienne kontrolne realizowanych procesów decyzyjnych*. Przyjęto, że są one zmiennymi symptomatycznymi podjętego ryzyka.

Hipoteza główna opracowania zakłada, że zmiany stanów ryzyka procesów decyzyjnych powodowane są zmianami uwarunkowań otoczenia społeczno-gospodarczego. W dowodzie hipotezy wykorzystano metody indukcyjne, w wyniku tego zostały zidentyfikowane zmienne kontrolne procesów decyzyjnych. Natomiast metody matematyczne, logiczne i statystyczne są instrumentami budowy modelu ryzyka oraz definicji miar ryzyka.

Celem opracowania jest uzasadnienie sformułowanej hipotezy badawczej. Dowód tezy wymaga przyjęcia dwóch istotnych założeń:

1. Zmiany uwarunkowań otoczenia są źródłem ryzyka procesów decyzyjnych.
2. Zmienne kontrolowane stanowią symptomatyczne oceny skutków ryzyka realizowanych procesów decyzyjnych.

Istotnym elementem dowodu sformułowanej hipotezy badawczej powinna być definicja stanu ryzyka. Stan ryzyka należy postrzegać jako zbiór statystycznych jego miar w określonym momencie czasu.

1. Uwarunkowania procesów decyzyjnych

Procesy decyzyjne związane z realizacją planów rozwoju organizacji gospodarczej mają cechy procesów warunkowych, są mianowicie determinowane uwarunkowaniami otoczenia, których z oczywistych powodów w procesach decyzyjnych nie można pominąć. Teoria zarządzania organizacją gospodarczą wyróżnia dwie rozłączne sfery oddziaływające na procesy decyzyjne: mikrootoczenie, w literaturze przedmiotu określane jako otoczenie branżowe oraz makrootoczenie krajowe i międzynarodowe. Funkcjonowanie organizacji jest możliwe pod warunkiem kształtowania właściwych relacji z otoczeniem tak, aby umożliwiały realizację takich działań, które nie naruszają istoty statusu jej planów strategicznych. Niezależnie od tego, czy jest to bliższe organizacji otoczenie – mikrootoczenie, czy dalsze otoczenie strategiczne – makrootoczenie, obydwie sfery oddziałują na organizację poprzez grupy interesów (*stake-*

* O znaczeniu zmiennych kontrolnych w procesach kontroli realizacji procesów decyzyjnych pisze T. Gołębiowski [4].

holders)*. Grupy interesów dysponują zasobami oraz kompetencjami, które mogą udostępniać organizacji bądź powstrzymać się od takich decyzji i działań. Skutki decyzji grup interesów mają istotne znaczenie dla organizacji, wpływają bowiem na pozycję konkurencyjną organizacji w sektorze, jej wyniki finansowe, a w krańcowo niekorzystnej sytuacji rynkowej decydują o jej przetrwaniu na rynku.

W mikrootoczeniu funkcjonuje wewnętrzna grupa interesów, reprezentowana przez menadżerów, pracowników oraz związki zawodowe funkcjonujące na poziomie organizacji gospodarczej. Zewnętrzne grupy interesów stanowią grupę finansującą organizację (właściciele, instytucje rynku kapitałowego, dostawcy, odbiorcy oraz konkurenci).

Uwarunkowania makrootoczenia są związane z interesem politycznym, społecznym i gospodarczym (administracja centralna i lokalna, organizacje polityczne, instytucje administracji centralnej, organizacje społeczne, związki zawodowe) [2, s. 37].

Identyfikacja grup interesów funkcjonujących w otoczeniu strategicznym organizacji jest niekiedy utrudniona. Pewne grupy funkcjonują dyskretnie, nie chcą ujawniać swego istnienia, ich działalność można oceniać jedynie po skutkach funkcjonowania grupy. Stosunkowo łatwo można zidentyfikować te grupy interesów, które funkcjonują w bliskim otoczeniu rynkowym organizacji. Są to grupy określane jako wewnętrzne oraz zewnętrzne, sygnalizują swoje oczekiwania i co istotne, istnieją metody szacowania siły ich wpływu **. Relacje organizacji z grupami interesów mogą mieć charakter kooperacji bądź konkurencyjny ***. Szczególne miejsce przypada grupom *stakeholders* z otoczenia międzynarodowego. Pomimo zacierania się granic pomiędzy otoczeniem krajowym i otoczeniem zagranicznym pewne grupy interesów są jednoznacznie umiejscowione w międzynarodowym otoczeniu strategicznym. Są to grupy konkurentów, dostawców surowców, materiałów, części, maszyn i urządzeń, technologii. W wyniku otwarcia gospodarek sprzyjających przepływowi kapitału oraz technologii istotnego znaczenia nabiera oddziaływanie grup interesów politycznych oraz społecznych.

* Termin ten oznacza interesariuszy, strategicznego „kibica” organizacji [7; 3].

** Proste, a jednocześnie skuteczne rozwiązanie proponuje K. Obłój [7, s. 118], dla określenia i zmierzenia siły oddziaływania grup interesów na organizację oraz siły relacji pomiędzy grupami otoczenia. Autor definiuje macierz, której elementami są wartości oznaczające siłę wpływu grup interesów. Inny problem to rozwiązanie kwestii pomiaru siły takich relacji, szczególnie trudne w przypadku oddziaływania grup, które nie ujawniają swoich interesów.

*** Klasycznym przykładem relacji konkurencyjnych są stosunki organizacji z jej konkurentami, relacje kooperacyjne nie występują w tak czystej postaci, są wynikiem kompromisów, ograniczających pragnienie dominacji jednej ze stron, uczestników umowy kooperacyjnej.

Struktura makrootoczenia jest umowna, gdyż nawet tak rozłączne dwa zbiory, jak otoczenie krajowe oraz międzynarodowe przenikają się, co oznacza, że otoczenie ekonomiczne ma zarówno cechy otoczenia krajowego, jak i międzynarodowego. Przyczyny tkwią w tendencjach do wzrostu współzależności gospodarek narodowych wskutek silnych tendencji integracyjnych czy globalizacji relacji politycznych i gospodarczych. Międzynarodowe wzorce konsumpcji, formy pracy i współpracy, styl spędzania wolnego czasu powodują, że czynniki społeczne i kulturowe stają się czynnikami nie tylko otoczenia krajowego. Silnie kojarzone z krajowym otoczeniem czynniki prawne i polityczne poprzez tworzenie ugrupowań regionalnych nabierają cech międzynarodowych. Nie popełniamy zatem błędu dzieląc jedynie formalnie makrootoczenie na krajowe i międzynarodowe.

Analizę makrootoczenia różnicują regionalny poziom rozwoju gospodarczego, czynniki kulturowe, polityka lokalnych władz administracji państwowej i samorządowej. Preferencje polityki lokalnej różnicują tempo rozwoju branż gospodarczych. Poprzez regulacje ekologiczne, rozwój lokalnych banków znających potrzeby ich klientów, zmieniając świadomość w sferze stylu życia lokalnych społeczności przesyłają czytelne sygnały w stronę organizacji produkujących „zdrową żywność”, sprzęt rekreacyjny, rozwijają usługi związane ze spędzaniem wolnego czasu.

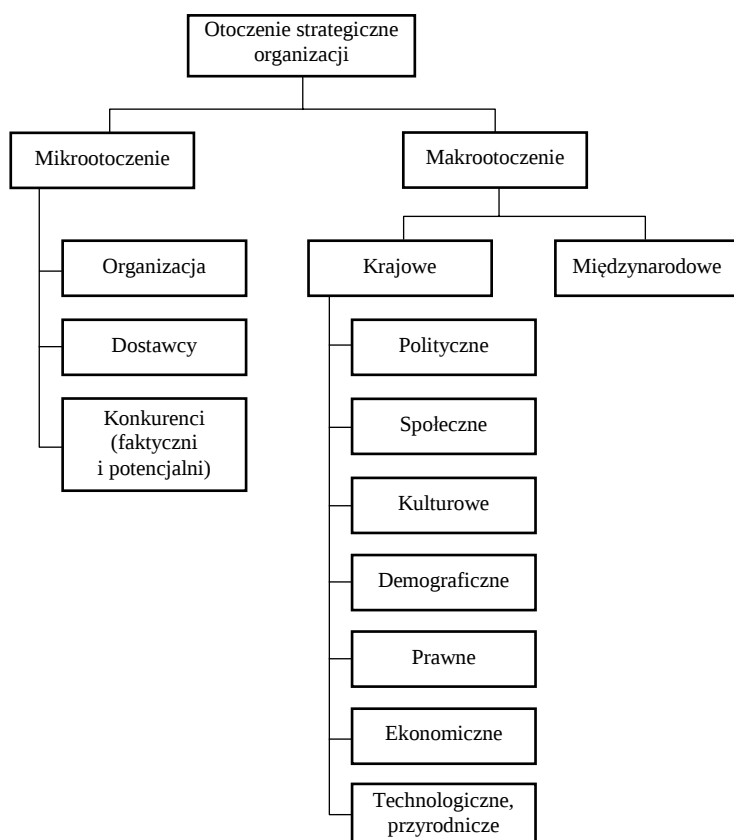
Analizując wpływ otoczenia na działalność organizacji trzeba mieć świadomość, że nie jest możliwe opracowanie jednolitej, wspólnej dla wszystkich organizacji analizy makrootoczenia, tak jak i to, że organizacja zwykle nie ma wpływu na sytuację w otoczeniu krajowym i zagranicznym*. Pomijając te utrudnienia analizy bezspornie można wyróżnić:

1. Otoczenie polityczne i prawne. Ma szczególnie istotne znaczenie dla realizacji wizji organizacji, zdefiniowanej w jej planach strategicznych, ponieważ od stopnia liberalizacji gospodarki zależy to, czy organizacje mogą się rozwijać w oparciu o kryteria ekonomiczne, czy ich rozwój zdeterminują zarządzenia, ograniczenia i zakazy pomijające reguły i prawa ekonomii. Charakteryzują je: krajowe regulacje prawne i administracyjne (normy prawa pracy, podatkowe, ochrony własności intelektualnej, normy ochrony środowiska, normy techniczne, normy międzynarodowe obowiązujące w kraju oraz funkcje państwa i organizacji społecznych)**.

* Wyjątek stanowią duże organizacje lub grupy organizacji, które są w stanie modyfikować stan otoczenia, definiując np. standardy technologiczne (Microsoft), trendy mody (Międzynarodowe Domy Mody) [4, s. 110 i nast.].

** Stopień oddziaływania administracji państwa na decyzje gospodarcze, udział własności państwa w gospodarce narodowej, zakres koncesjonowania działalności gospodarczej, zakres i formy preferencji, np. formy ulg podatkowych, zamówień publicznych, subwencji, zakres i formy działań organizacji finansowanych przez budżet państwa, wspierających działalność gospodarczą.

2. Otoczenie społeczne i kulturowe. Poprzez stopień sprawności sfery usług publicznych, świadomość ekologiczną, wzorce zachowań rynkowych, tendencje w zmianach organizacji czasu wolnego, zmiany oczekiwań pracowników wobec pracodawców zostają zdefiniowane cechy otoczenia społeczno-kulturowego. Są nimi: poziom cywilizacyjny, określone postawy społeczne oraz wzorce konsumpcji.



Rys. 1. Otoczenie strategiczne organizacji

Źródło: Opracowanie własne na podstawie [1, s. 102].

3. Otoczenie ekonomiczne. Definiuje powszechne w świadomości społecznej cechy, które obok zbioru miar koniunktury identyfikują instytucje organizujące i wspierające działalność gospodarczą. Ponadto są to cechy opisujące zmiany strukturalne w gospodarce kraju. Właściwością otoczenia ekonomicznego jest zbiór miar, które je opisuje. Ma to szczególne znaczenia wówczas, gdy zachodzi konieczność precyzyjnego opisu zachodzących w otoczeniu zmian. Miary zmian stanu otoczenia mogą mieć cechy ilościowe bądź

jakościowe. Właściwości miar mają istotne znaczenie w ocenie zmian uwarunkowań otoczenia. Dużą pojemność informacyjną mają miary ilościowe, niestety takiego „komfortu” nie dają cechy jakościowe, które mają wartość subiektywną (np. organizacja instytucji infrastruktury gospodarczej). Otoczenie ekonomiczne charakteryzują:

- wielkość i tempo wzrostu PKB, poziom i zmiany stopy procentowej, zmiany kursów walut, poziom inflacji, stopa konsumpcji, oszczędności, poziom zadłużenia wewnętrznego, zagranicznego, poziom bezrobocia, wydajności pracy, produktywności zasobów, nakłady inwestycyjne, amortyzacja majątku produktywnego, polityka monetarna i fiskalna państwa,
 - organizacja oraz sprawność funkcjonowania instytucji infrastruktury gospodarczej,
 - tendencje i tempo zmian strukturalnych*.
4. Otoczenie technologiczne, a także przyrodnicze. W obecnej dobie stanowi najbardziej istotne zagrożenia dla funkcjonujących organizacji, ale jest także znaczącą szansą ich rozwoju. Do istotnych czynników tego otoczenia należy zaliczyć:
- kierunki zmian postępu technicznego, które identyfikują dostępność nowoczesnych technologii, cykle życia produktów oraz technologii,
 - jakość techniczną substytutów,
 - skalę i obszar działań instytucji prowadzących badania nad nowymi technologiami [4, s. 116],
 - stan i dostępność zasobów naturalnych,
 - czynniki naturalne mające wpływ na koszty funkcjonowania organizacji,
 - zmiany stanu środowiska naturalnego.
5. Otoczenie demograficzne. Determinuje założenia rozwoju społeczno-gospodarczego kraju. W plany tych założeń wpisują się organizacje gospodarcze definiujące swoje własne strategie rozwoju, uwzględniające potrzeby zmieniającego się rynku pod wpływem czynników demograficznych. Otoczenie to charakteryzują:
- liczba ludności kraju lub regionu działania organizacji,
 - struktura oraz gęstość zaludnienia,
 - przyrost naturalny i struktura wiekowa ludności,
 - struktura narodowościowa (etniczna),
 - intensywność i kierunki migracji,
 - poziom aktywności zawodowej,

* Tendencje zmian strukturalnych są często niedoceniane, zaniechanie monitorowania tej cechy otoczenia ekonomicznego uniemożliwia śledzenie zagrożeń na rynkach „schyłkowych”. Systematyczny monitoring tendencji oraz tempa zmian pozwala identyfikować rynki „niszowe”, w odpowiednim momencie przejąć zasoby pracownicze z organizacji, które w wyniku zmian strukturalnych musiały ograniczyć działalność, zmienić profil działalności bądź „znikły” z rynku.

- struktura zawodowa,
- wielkość i struktura gospodarstw domowych,
- rozkład dochodów w grupach ludności [4, s. 113].

Czynniki demograficzne wyznaczają skalę oraz strukturę popytu. Zróżnicowany poziom dochodów decyduje o strukturze asortymentowej produktów, determinowanej różnym poziomem jakości i cen. Natężenie ruchów migracyjnych zmienia strukturę rynku pracy, różnicując koszty pracy. Natomiast wzrost aktywności zawodowej powoduje rosnący popyt na usługi związane z opieką nad dziećmi. Rosną usługi na przetworzoną żywność.

Otoczeniem organizacji postrzeganym z perspektywy mikro jest sektor (branża), w którym podmiot gospodarczy funkcjonuje. Sektor jest zbiorem organizacji wytwarzających identyczne bądź podobne produkty-substytuty. Substytucyjność produktów sprawia, że pojęcie „sektora” nie jest jednoznacznie określone*. Większą dokładność identyfikacji sektora uzyskamy charakteryzując jego cechy. Są nimi:

- wielkość rynku, mierzona skalą popytu na produkty sektora,
- tempo przyrostu popytu w sektorze,
- zasięg konkurencji w sektorze (lokalny, krajowy, międzynarodowy),
- stopień koncentracji podaży,
- stopień koncentracji popytu,
- zakres oraz możliwości integracji pionowej organizacji w sektorze,
- poziom barier wejścia i wyjścia w sektorze,
- tempo zmian technologii,
- dominujący rodzaj kanałów dystrybucji,
- stopień zróżnicowania produktów konkurujących organizacji sektora,
- intensywność korzyści skali w produkcji, logistyce i marketingu,
- poziom zależności kosztów produktów od stopnia wykorzystania potencjału,
- efekt doświadczenia a skala działania,
- rentowność sektora w porównaniu z innymi branżami w kraju (tego samego sektora w innym kraju).

Cechy sektora nie są jedynymi cechami definiującymi mikrootoczenie, obok wymienionych charakterystyczny dla tego bezpośredniego otoczenia organizacji gospodarczej jest zbiór interesariuszy i związane z tym uwarunkowania dotyczące dostawców oraz odbiorców. Ponadto zidentyfikowany zbiór uwarunkowań uzupełniają te, które identyfikują relacje z konkurencją organizacji na rynku sektora [8, s. 282 i nast.].

Uwarunkowania otoczenia organizacji gospodarczej są w większości mierzalne. Odrębne zagadnienie stanowią uwarunkowania jakościowe, w takim przypadku stajemy przed zagadnieniem pomiaru tego rodzaju uwarunkowań.

* Na przykład sektor zabawek dla dzieci jest obszerny, gry komputerowe zawężają obszar sektora.

Pominięcie tak istotnego nośnika informacji o dynamice zmian otoczenia byłoby błędem metodologicznym. Jak rozwiązać zagadnienie pomiaru uwarunkowań niemierzalnych, aby w oparciu o nie mieć możliwość definicji zmiennych kontrolnych? Teoretycznie, a także jak potwierdza to praktyka, problem ten można rozwiązać definiując wzorce (normy) uwarunkowań. Odchylenie pomiędzy uwarunkowaniem w momencie $t \in [1, T]$ a jego wzorcem można definiować ilościowo, co powinno rozwiązać problem aproksymacji podzbioru uwarunkowań niemierzalnych zmiennymi mierzalnymi.

2. Zarządzanie w warunkach ryzyka

Zarządzanie w warunkach podjętego ryzyka jest procesem realizacji celów przyjętych w planach rozwoju organizacji gospodarczej, ale jest procesem warunkowym. Warunkowość procesu wynika z konieczności uwzględniania zmienności uwarunkowań otoczenia organizacji gospodarczej. Ryzyko realizowanych procesów decyzyjnych jest związane z dynamiką zmienności uwarunkowań otoczenia, a bardziej precyzyjnie z nieprzewidywalnym kierunkiem i tempem zmian istotnych dla realizowanych celów uwarunkowań. Monitorowanie zmian uwarunkowań otoczenia jest integralnym elementem procesów zarządzania organizacją gospodarczą i stanowi strukturalną składową funkcji kontroli realizowanych procesów decyzyjnych.

Zarządzanie w warunkach ryzyka jest łączeniem procesów decyzyjnych z instrumentami ochrony realizowanych celów przed skutkami podjętego ryzyka. Rodzaj instrumentów, czas pozyskania, ich efektywność, a także koszty realizacji procesów zarządzania definiują kryterium wyboru koniecznych zabezpieczeń przed skutkami ryzyka. Istotny element procesu zarządzania stanowi moment uruchomienia instrumentów zabezpieczeń. Jest on uzależniony od oceny stopnia zagrożenia realizowanych w procesie decyzyjnym celów. Jakość oraz wiarygodność ocen stanu ryzyka w każdym momencie procesu decyzyjnego ma istotny wpływ na podejmowane decyzje. Może powodować zmiany tylko niektórych elementów decyzji bądź też w skrajnie niekorzystnych przypadkach konieczność przerwania procesu decyzyjnego.

Monitorowanie zmian uwarunkowań otoczenia, obok informacji o dynamice i kierunku zmian przyjętych w planach założeń realizacji celów, może i powinno być wykorzystane do konstrukcji instrumentów ostrzegania o zmianach stanów ryzyka.

Wspomaganie procesów decyzyjnych wykorzystujące instrumenty ostrzegawcze, mające za założenia informować o zmianach tendencji oraz dynamiki stanów ryzyka wymaga budowy bazy miar „stanów bazowych ryzyka”^{*}.

3. Pomiar ryzyka. Stan ryzyka w procesach zarządzania

W pracy [11, s. 91 i nast.] przedstawiono konstrukcję modelu ryzyka. Model jest wektorem losowym o składowych będących zmiennymi kontrolnymi ryzyka realizowanego procesu decyzyjnego. Przyjmując założenie o rodzaju rozkładu funkcji gęstości prawdopodobieństwa, dla oceny miar ryzyka można zdefiniować jego statystyczne miary:

- prawdopodobieństwo tego, że składowe wektora ryzyka przyjmą wartości z określonego przedziału zmienności,
- wartość oczekiwana wektora losowego,
- wariancja wektora losowego,
- kowariancja wektora losowego (macierz kowariancji pomiędzy składowymi wektora ryzyka)^{**}.

Stan ryzyka procesu decyzyjnego określają jego miary w określonym momencie realizacji przyjętych w planach organizacji gospodarczej celów. Stan ryzyka jest więc zbiorem czterech miar statystycznych wektora ryzyka:

$$SR(P(X), E(X), Var(X), Cov(X)) \quad (1)$$

Miary stanu ryzyka nie są typologicznie jednorodne, prawdopodobieństwo zdarzenia, że składowe wektora ryzyka przyjmą wartości z określonych przedziałów zmienności – $P(X)$ jest skalar, wartość oczekiwana oraz wariancja wektora ryzyka – $E(X), Var(X)$ są wektorami, natomiast kowariancje składowych wektora ryzyka – $Cov(X) = cov(X_i, X_j)$, gdzie $i, j = 1, 2, \dots, k$ oraz $i \neq j$ definiują macierz kwadratową $k - 1$ wymiarową.

^{*} Stan bazowy definiują uwarunkowania otoczenia przyjęte w założeniach planu rozwoju organizacji gospodarczej, które umożliwiają realizację przyjętych w planach celów.

^{**} Niech $X = (X_1, X_2, \dots, X_k)$ jest wektorem ryzyka identyfikowanym ze zbiorem $\{X_1, X_2, \dots, X_k\}$ zmiennych kontrolnych realizowanego procesu decyzyjnego. Przyjmując hipotezę, że funkcja rozkładu gęstości prawdopodobieństwa f ma rozkład logarytmiczno-normalny definiujemy miary statystyczne wektora:

1. Prawdopodobieństwo zdarzeń, że zmienne losowe X_h , gdzie: $h = 1, 2, \dots, k$ przyjmą wartości x_h należące do przedziału $[a_h, b_h]$: $P(a_1 \leq X_1 \leq b_1, \dots, a_k \leq X_k \leq b_k)$.
2. Wartość oczekiwana wektora ryzyka: $E(X) = (E(X_1), \dots, E(X_k))$.
3. Wariancja wektora ryzyka: $Var(X) = (Var(X_1), \dots, Var(X_k))$.
4. Kowariancje składowych wektora ryzyka: $Cov(X_i, X_j)$ gdzie: $i, j = 1, \dots, k$ oraz $i \neq j$.

Założmy, że oszacowano stan ryzyka w przedziale czasu $[t, t+1]$, oznaczając, że dysponujemy informacją $SR^{(t)}$ oraz $SR^{(t+1)}$. Relacja pomiędzy $SR^{(t)}$ oraz $SR^{(t+1)}$ stanowi istotną informację w ocenie kierunku oraz dynamiki zmian ryzyka towarzyszącego realizacji procesów decyzyjnych. Co oznaczają te zmiany w przypadku, gdy składowe stanu ryzyka są tak zróżnicowane typologicznie. W przypadku pierwszej składowej stanu ryzyka różnica $(P^{(t+1)} - P^{(t)})$ pozwala określić kierunek zmian zmiennych kontrolnych, z jaką przyjmują wartości z określonych przedziałów zmienności w czasie $[t+1, t]$. Składowe druga i trzecia stanu ryzyka są wektorami, stąd miarą zmian stanu ryzyka w przedziale $[t+1, t]$ są np. zmiany odległości pomiędzy wektorami: $d(E^{(t+1)}, E^{(t)})$, $d(Var^{(t+1)}, Var^{(t)})^*$.

Czy badanie zmian przy wykorzystaniu miar definiujących zmiany w kolejno następujących po sobie momentach $[t+1, t]$ kreśli właściwy z pozycji zarządzających obraz zmian stanów ryzyka realizowanych procesów decyzyjnych? Dla oceny wartości zdefiniowanych miar należy przywołać pojęcie „ryzyka procesu decyzyjnego” – wyraża się ono różnicą pomiędzy oczekiwaniami zapisanymi w planach rozwoju organizacji a realizacją planów. Ta interpretacja miary stanu ryzyka odwołująca się do „stanów bazowych ryzyka” zdefiniowanych w planach organizacji sugeruje dokonanie istotnej modyfikacji miary zmian stanów ryzyka.

Niech $SR^{(b)}$, tzn. $(P^{(b)}(X), E^{(b)}(X), Var^{(b)}(X), Cov^{(b)}(X))$ oznacza „stan bazowy”, należy zmodyfikować definicje zmian stanów ryzyka, są nimi różnica $(P^{(t)} - P^{(b)})$ oraz odległości: $d(E^{(t)}, E^{(b)})$, $d(Var^{(t)}, Var^{(b)})$. Modyfikacja miar zmienia interpretację zmian stanów ryzyka^{**}. W każdym momencie t pokazują zmiany pomiędzy miarami składowymi stanu ryzyka w relacji do stanu zdefiniowanego w planie. Pokazują zmiany $SR^{(t)}$ w relacji do $SR^{(b)}$, co bardziej wiarygodnie mierzy zmienność stanu ryzyka w czasie.

Źródłem zmienności stanów ryzyka są zmiany uwarunkowań otoczenia organizacji gospodarczej realizującej swoje plany rozwoju gospodarczego, składając się na zmiany zmiennych kontrolnych $\{X_h\}$, $h=1,2,...,k$. Zbiór tych wyróżnionych zmiennych identyfikuje ryzyko realizowanego procesu decyzyjnego.

* Dane dwa wektory: $X = (x_1, x_2, ..., x_n)$, $Y = (y_1, y_2, ..., y_n)$, gdzie $X, Y \in R^n$ oraz pewna symetryczna, dodatnio określona macierz C . Miara $d(X, Y) = \sqrt{(X - Y) C^{-1} (X - Y)^T}$ jest odległością pomiędzy wektorami $X, Y \in R^n$ w sensie Mahalanobisa [6].

** Pojęcie „stan bazowy” jest określeniem sformułowanym w tym opracowaniu i dla jego potrzeb. Oznacza przyjęte w planach założenia ilościowe bądź jakościowe o zmiennych kontrolnych planowanych procesów decyzyjnych.

Przyjmijmy, że k -wymiarowy wektor losowy $X = (X_1, X_2, \dots, X_k)$ jest modelem ryzyka i jego rozkład jest zgodny z rozkładem normalnym*. Funkcja gęstości wektora X o wektorze wartości oczekiwanych $E(X) = (E(X_1), E(X_2), \dots, E(X_k))$ i macierzy kowariancji $\Sigma = [\text{cov}(X_i, X_j)]$, gdzie $i, j = 1, 2, \dots, k$ oraz $i \neq j$ jest równa:

$$f_{E(X), \Sigma} = \frac{1}{(2\pi)^{k/2} |\Sigma|^{1/2}} \exp \left(-\frac{1}{2} (X - E(X)) \Sigma^{-1} (X - E(X))^T \right) \quad (2)$$

Wykładnik licznika związku (2) jest kwadratem odległości w sensie Mahalanobisa pomiędzy dwoma wektorami w przestrzeni k -wymiarowej**. Odległość w sensie Mahalanobisa opisuje właściwości pewnego zbioru, w analizowanym przypadku zmian odległości pomiędzy składowymi stanu ryzyka (drugą i trzecią, tj. $(E^{(t)}(X), \text{Var}^{(t)}(X))$) a ich odpowiednimi miarami „stanu bazowego” $(E^{(b)}(X), \text{Var}^{(b)}(X))$. Odpowiedź na pytanie, czy analiza składowych stanu ryzyka pozwala uznać, że należą jeszcze do klasy, która w stanie ryzyka w momencie t nie dostrzega zagrożeń realizowanego procesu decyzyjnego wiąże się z odpowiedzią na pytanie, czy skupienie miar stanu ryzyka o określonym momencie w relacji do miar „stanu bazowego” jest wystarczająco bliskie. Skupienie to jest obrazem podobieństwa wektora $E^{(t)}(X)$ do $E^{(b)}(X)$ (bądź $\text{Var}^{(t)}(X)$ do $\text{Var}^{(b)}(X)$)***.

Oznaczenie stanu ryzyka identyfikowanego z otoczeniem pojmowanym jako „bezpiecznie bliskie” wiąże się z odpowiedzią na pytanie o wymiar odległości pomiędzy miarami stanu ryzyka w momencie t , tzn. $(E^{(t)}(X), \text{Var}^{(t)}(X))$, a $(E^{(b)}(X), \text{Var}^{(b)}(X))$. Punkty o identycznej odległości od punktu centralnego (wyznaczają go miary „stanu bazowego” – $(E^{(b)}(X), \text{Var}^{(b)}(X))$) tworzą w przestrzeni $k+1$ wymiarowej elipsoidę hipersferyczną. Zachodzą przy tym trzy szczególne przypadki:

1. Składowe wektora ryzyka mają identyczne wariancje (po standaryzacji można przyjąć, że są równe 1), nie są pomiędzy sobą skorelowane, oznacza to, że macierz kowariancji jest macierzą jednostkową (miara w sensie Mahalanobisa jest równa odległości w sensie Euklidesa), elipsoida hipersferyczna

* Pod warunkiem, że dowolna kombinacja liniowa $Y = \alpha_1 X_1 + \alpha_2 X_2 + \dots + \alpha_k X_k$ składowych wektora X jest także zgodna z rozkładem normalnym ($\alpha_i \in R$, $i = 1, 2, \dots, k$).

** Ta istotna dla prowadzonego wywodu zgodność eksponuje znaczenie macierzy kowariancji Σ w ocenie zmian stanu ryzyka w przedziale czasu $[t+1, t]$.

*** Konstrukcję miary podobieństwa oparto na informacji o wariancjach składowych wektora ryzyka oraz korelacjach pomiędzy nimi w momencie t .

redukuje się w tym przypadku do sfery o środku w punkcie $(E^{(b)}(X_1), E^{(b)}(X_2), \dots, E^{(b)}(X_k)) / (Var^{(b)}(X_1), Var^{(b)}(X_2), \dots, Var^{(b)}(X_k)) /$ o równaniu:

$$d(E(X)^{(t)}, E(X)^{(b)}) = \sqrt{(E^{(t)}(X) - E(X)^{(b)}(X)) I^{-1} (E^{(t)}(X) - E(X)^{(b)})^T}$$

bądź w przypadku trzeciej składowej stanu ryzyka:

$$d(Var(X)^{(t)}, Var(X)^{(b)}) = \sqrt{(Var^{(t)}(X) - Var(X)^{(b)}(X)) I^{-1} (Var^{(t)}(X) - Var(X)^{(b)})^T}$$

2. Składowe wektora ryzyka nie są skorelowane, mają różne wariancje $(\sigma_1^2, \sigma_2^2, \dots, \sigma_k^2)$, macierz kowariancji D jest macierzą diagonalną, obrazem relacji jest elipsoida hipersferyczna o środku w punkcie $(E^{(b)}(X_1), E^{(b)}(X_2), \dots, E^{(b)}(X_k)) / (Var^{(b)}(X_1), Var^{(b)}(X_2), \dots, Var^{(b)}(X_k)) /$ o równaniu:

$$d(E(X)^{(t)}, E(X)^{(b)}) = \sqrt{(E^{(t)}(X) - E(X)^{(b)}(X)) D^{-1} (E^{(t)}(X) - E(X)^{(b)})^T} \quad \text{bądź}$$

$$d(Var(X)^{(t)}, Var(X)^{(b)}) = \sqrt{(Var^{(t)}(X) - Var(X)^{(b)}(X)) D^{-1} (Var^{(t)}(X) - Var(X)^{(b)})^T}$$

Pozwala to zdefiniować odległość:

$$d(E(X)^{(t)}, E(X)^{(b)}) = \sqrt{\left(\frac{(E^{(t)}(X_1) - E^{(b)}(X_1))^2}{\sigma_1^2} + \dots + \frac{(E^{(t)}(X_k) - E^{(b)}(X_k))^2}{\sigma_k^2} \right)} \quad \text{bądź}$$

$$d(Var(X)^{(t)}, Var(X)^{(b)}) = \sqrt{\left(\frac{(Var^{(t)}(X_1) - Var^{(b)}(X_1))^2}{\sigma_1^2} + \dots + \frac{(Var^{(t)}(X_k) - Var^{(b)}(X_k))^2}{\sigma_k^2} \right)}$$

3. Składowe wektora ryzyka są skorelowane, mają różne wariancje $(\sigma_1^2, \sigma_2^2, \dots, \sigma_k^2)$, macierz kowariancji C nie jest macierzą diagonalną, obrazem relacji jest elipsoida hipersferyczna o środku w punkcie $(E^{(b)}(X_1), E^{(b)}(X_2), \dots, E^{(b)}(X_k)) / (Var^{(b)}(X_1), Var^{(b)}(X_2), \dots, Var^{(b)}(X_k)) /$ o równaniu:

$$d(E(X)^{(t)}, E(X)^{(b)}) = \sqrt{(E^{(t)}(X) - E(X)^{(b)}(X)) C^{-1} (E^{(t)}(X) - E(X)^{(b)})^T}$$

bądź w przypadku trzeciej składowej stanu ryzyka:

$$d(Var(X)^{(t)}, Var(X)^{(b)}) = \sqrt{(Var^{(t)}(X) - Var(X)^{(b)}(X)) C^{-1} (Var^{(t)}(X) - Var(X)^{(b)})^T}$$

Obraz elipsoidy hipersferycznej jest obrócony o kąt wyznaczony przez macierz wektorów własnych macierzy C , długości osi elipsoidy odpowiadają pierwiastkom kwadratowym pierwiastków własnych $(\lambda_1^2, \lambda_2^2, \dots, \lambda_k^2)$ macierzy C .

W definicji stanu ryzyka konstrukcja wszystkich jego składowych jest oparta na elementach macierzy wariancji i kowariancji składowych wektora ryzyka. Zmiany elementów macierzy informują o zmianie stanu ryzyka. To istotna informacja, która jeśli proces monitorowania otoczenia prowadzony jest systematycznie z określoną częstotliwością pozwala określać ten moment krytyczny, w którym zmieniają się parametry stanu ryzyka. Zmiana może sprzyjać realizacji celów zdefiniowanych w planach rozwoju organizacji gospodarczej bądź przeciwnie. Niezależnie od kierunku zmian, takie informacje umożliwiają formułowanie odpowiedzi na wątpliwości związane ze stopniem zagrożenia realizowanych procesów decyzyjnych, co będzie miało istotne znaczenie w doborze instrumentów zabezpieczeń przed skutkami ryzyka.

Zakończenie

Idea budowy instrumentów ostrzegania o zmianie stanów ryzyka w wyniku zmian zachodzących w otoczeniu społeczno-gospodarczym, istotnego narzędzia systemu wspomagania procesów decyzyjnych realizowanych w warunkach ryzyka, ma wymiar teoretyczny. Perspektywy uczynienia z nich skutecznego i efektywnego narzędzia praktyki zarządzania są obiecujące.

Zarządzanie jest procesem realizacji planów. Poprawnie sporządzony plan powinien uwzględniać te wszystkie uwarunkowania otoczenia, które będą miały wpływ na wynik – cel jego realizacji. Oznacza to, że zarządzający dysponują zbiorem miar „stanów bazowych”, a w wyniku prowadzonego procesu monitoringu zmian uwarunkowań otoczenia także miarami stanów ryzyka w każdym momencie określonym w założeniach monitoringu. Porównanie miar „stanów bazowych” – wzorca stanu ryzyka z miarami stanu w monitorowanym momencie wspiera procesy decyzyjne w tych momentach krytycznych realizowanego procesu zarządzania, które wymagają podjęcia decyzji o rodzaju instrumentów zabezpieczających przed skutkami ryzyka.

Literatura

1. Dawid F.R., *Concepts of Strategic Mangement*, Prentice Hall, Upper Sadle River, New York 1997.
2. Donaldson G., Lorsch J.W., *Decion Making At The Top: The Shaping of Strategic Direction*, Basic Books, New York 1983.
3. Freeman R.E., *Strategic Management*, Pitman, Boston 1984.
4. Gołębiowski T., *Zarządzanie strategiczne. Planowanie i kontrola*, Difin, Warszawa 2001.

5. McMillan I.C. Jones P.E., *Strategy Formulation: Power and Politics*, West Publishing, St. Paul 1986.
6. Mahalanobis P.C., *On the generalised distance in statistics*, Proceedings of the National Institute of Sciences of India 2, 1936
7. Obłój K., *Strategia organizacji*, PWE, Warszawa 1999.
8. Rokita J., *Zarządzanie strategiczne. Tworzenie i utrzymanie przewagi strategicznej*, PWE, Warszawa 2005.
9. Willet A.H., *The Economic Theory of Risk Insurance*, University of Pennsylvania Press, Philadelphia 1951.
10. *Zarządzanie ryzykiem*, red. K. Jajuga, PWN, Warszawa 2007.
11. Zemke J., *Ryzyka zarządzania organizacją gospodarczą*, Uniwersytet Gdański, Gdańsk 2009.

RISK MANAGEMENT IN THE CONTEXT OF IN A DIVERSE ENVIRONMENT OF SOCIO – ECONOMIC

Summary

The content of paper is limited to provide sources of risk, issues of risk identification processes of decision-making and risk measurement. The development occupies a special place the issue of sources of risk and its identification. Economic organization management theory draws attention to the so-called. control variables of the processes of decision-making. The study assumes that they are symptomatic of the chance variables. Home to develop a hypothesis assumes the risk that changes in state decision-making processes are caused by changes in the socio economic conditions.

Maria Balcerowicz-Szcutnik

Uniwersytet Ekonomiczny w Katowicach

UWARUNKOWANIA POZIOMU BEZROBOCIA WYBRANYCH PAŃSTW UE – – ANALIZA STATYSTYCZNA

Wprowadzenie

Bezrobocie zawsze było jednym z najważniejszych problemów gospodarczo-społecznych w Unii Europejskiej, a jego zwalczanie stawało się priorytetem dla kolejnych ekip rządzących poszczególnych państw unijnych. Niski poziom bezrobocia to podstawowy warunek umożliwiający utrzymanie i rozwijanie dobrobytu, natomiast wysoki poziom braku zatrudnienia ma określone konsekwencje gospodarcze i społeczne. Wśród gospodarczych skutków bezrobocia wystarczy wymienić potrzebę finansowania kosztów bezrobocia i jego rozlicznych skutków, a także konsekwencje niepełnego wykorzystania zasobów pracy i potencjału ludzkiego. Wśród społecznych skutków bezrobocia za najbardziej dotkliwe można uznać spadek dochodów ludności, rozszerzanie się kręgów ubóstwa oraz degradację psychiczną i moralną osób niemających pracy. W latach 2008 i 2009 można było zaobserwować gwałtowny wzrost stopy bezrobocia, co było swoistym echem światowego kryzysu gospodarczego. W połowie 2008 r. światowa gospodarka znalazła się w najgłębszej recesji od czasów drugiej wojny światowej. Fazę dekoniunktury pogłębił krach na rynkach finansowych zapoczątkowany w Stanach Zjednoczonych. Spadek PKB w niemal wszystkich krajach na świecie nie pozostał bez echa na rynkach pracy. Pierwsze skutki światowego kryzysu gospodarczego w postaci wzrastającego bezrobocia pojawiły się w Europie już na początku 2008 r. W całej Unii Europejskiej liczba osób pozostających bez pracy zaczęła systematycznie wzrastać od czerwca 2008 r. Wówczas stopa bezrobocia kształtowała się na poziomie 6,9%. W marcu 2009 r. w całej UE wskaźnik ten wynosił już 8,3%*. Pomiędzy styczniem 2008 r. a marcem 2009 r. Wspólnocie przybyło 3,9 mln osób bezrobotnych, co było wyraźnym sygnałem gwałtownie pogarszającej się kondycji unijnego rynku pracy.

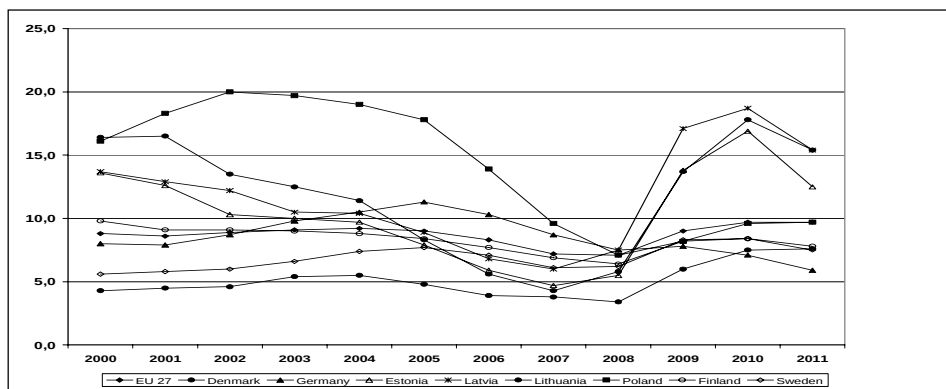
* Dane Eurostat.

Prezentowane opracowanie zawiera analizę podstawowego miernika poziomu bezrobocia, a mianowicie analizę zharmonizowanej stopy bezrobocia dla wybranej grupy państw nadbałtyckich oraz stopy bezrobocia długoterminowego*. Wybór państw objętych analizą jest nieprzypadkowy. Cztery z nich, czyli Niemcy, Dania, Szwecja i Finlandia to państwa tzw. starej unii, a cztery pozostałe to nowo przyjęte do UE państwa nadbałtyckie, czyli Polska, Litwa, Łotwa i Estonia. Zwłaszcza te trzy ostatnie powinny być przedmiotem szczególnej analizy, ponieważ jeszcze w 2007 r. Litwę, Łotwę i Estonię określano mianem „europejskich tygrysów”. Gospodarki w tych krajach rozwijały się najszybciej spośród wszystkich państw należących do Wspólnoty. W 2008 r. weszły natomiast w fazę największej recesji od czasów uzyskania przez nie niepodległości na początku lat 90. XX w. Litwa, Łotwa i Estonia to państwa bardzo mocno uzależnione od światowej koniunktury, dlatego pierwsze oznaki globalnego kryzysu gospodarczego pojawiły się właśnie tam. Do recesji wymienionych krajów bałtyckich doprowadziły także zmniejszenie się popytu wewnętrznego, spadek produkcji przemysłowej, ograniczenie napływu bezpośrednich inwestycji zagranicznych oraz wzrost inflacji. Wraz z dekonunkturą nastąpił bardzo szybki przyrost liczby osób bezrobotnych. Od połowy 2008 r. Litwa, Łotwa i Estonia odznaczają się największą dynamiką wzrostu bezrobocia wśród wszystkich państw UE. Przedstawione analizy szczegółowe można podzielić na dwie zasadnicze części. Pierwsza będzie dotyczyć analizy dynamiki stopy bezrobocia ogółem za lata 2000-2011 w ujęciu rocznym, a druga analizy bezrobocia długookresowego, szczególnie dotkliwego dla samych zainteresowanych i kłopotliwego ze względu na metody łagodzenia jego skutków.

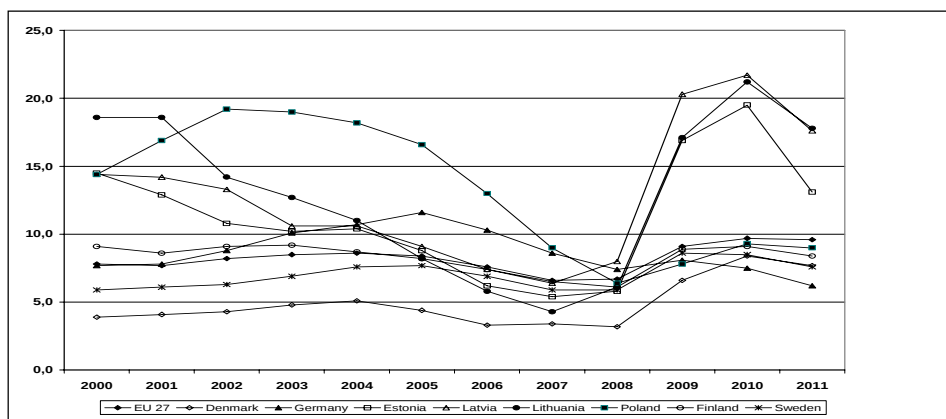
1. Bezrobotni jako szczególna kategoria na rynku pracy

Oceny poziomu bezrobocia dokonamy według odpowiednich wskaźników statystycznych, czyli miernika stopy bezrobocia ogółem oraz z uwzględnieniem płci pracownika. Wykresy na rys. 1, 2, 3 przedstawiają empiryczne linie trendu stopy bezrobocia w latach 2000-2011 w wymienionych krajach nadbałtyckich UE w ujęciu ogółem i z uwzględnieniem płci potencjalnego pracownika.

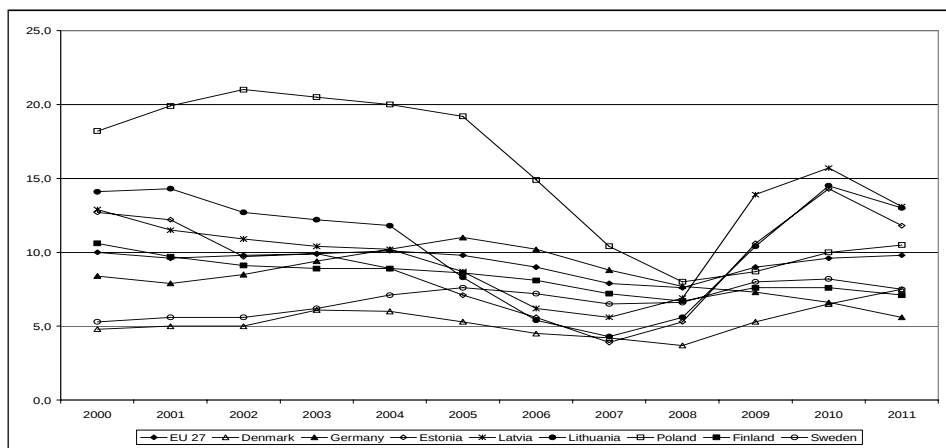
* Eurostat mierzy zharmonizowaną stopę bezrobocia jako procent osób w wieku 15-74 lata pozostających bez pracy, zdolnych podjąć zatrudnienie w ciągu najbliższych dwóch tygodni, którzy aktywnie poszukiwali pracy w ciągu ostatnich tygodni w odniesieniu do wszystkich osób aktywnych zawodowo w danym kraju.



Rys. 1. Dynamika zharmonizowanej stopy bezrobocia w państwach nadbałtyckich za lata 2000-2011 (ogółem)



Rys. 2. Dynamika zharmonizowanej stopy bezrobocia w państwach nadbałtyckich za lata 2000-2011 (mężczyźni)



Rys. 3. Dynamika zharmonizowanej stopy bezrobocia w państwach nadbałtyckich za lata 2000-2011 (kobiety)

Szczegółowa analiza powyższych linii trendu pozwala na stwierdzenie, że:

- Dla ogółu państw UE stopa bezrobocia w ostatniej dekadzie uległa nieznacznemu spadkowi do 2008 r. zarówno ogółem, jak i z uwzględnieniem płci pracownika. Po 2008 r., jak wcześniej stwierdzono, poziom bezrobocia gwałtownie rośnie, co niezaprzeczalnie jest echem światowego kryzysu gospodarczego. Jest to widoczne zwłaszcza w latach 2008-2010. W 2011 r. pojawia się nieznaczny spadek poziomu bezrobocia, lecz nie można jeszcze stwierdzić, czy ta tendencja utrzyma się czy ulegnie zmianie. Prognostycy banków światowych są przekonani, że zjawisko bezrobocia będzie się nasilać i wystąpi swoiste echo kryzysu z 2008 r.
- W państwach skandynawskich poziom bezrobocia jest niższy niż w całej Unii Europejskiej, co nie jest zaskoczeniem z uwagi na fakt, że współczynniki aktywności zawodowej w tych państwach miały zdecydowanie wyższy poziom niż w całej UE.
- Również państwa skandynawskie nie ustrzegły się przed gwałtownym wzrostem poziomu bezrobocia po 2008 r. Wprawdzie poziom bezrobocia był tam niższy niż w innych państwach unijnych, ale efekt „odbicia” spowodowany krachem finansowym był widoczny.
- Warto zauważyć zdecydowaną rozbieżność stopy bezrobocia w analizowanych państwach nadbałtyckich, np. najwyższa stopa bezrobocia w Danii lub Szwecji w całym okresie 2000-2011 wynosząca około 8% była równa najniższej w Polsce za wskazany okres.

- Te prawidłowości dynamiki utrzymują się również w odniesieniu do płci osoby pozostającej bez pracy. Różnice widoczne są jedynie w wartościach mierników z oczywistą przewagą na niekorzyść kobiet.

W tab. 1 zawarto wartości współczynników określających stopę bezrobocia w latach 2000, 2005, 2011.

Tabela 1

Stopa bezrobocia w państwach nadbałtyckich w latach 2000, 2005, 2011,
z uwzględnieniem płci

	2000 r.			2005 r.			2011 r.		
	ogółem	mężczyźni	kobiety	ogółem	mężczyźni	kobiety	ogółem	mężczyźni	kobiety
EU 27	8,8	7,8	10,0	9,0	8,4	9,8	9,7	9,6	9,8
Denmark	4,3	3,9	4,8	4,8	4,4	5,3	7,6	7,7	7,5
Germany	8,0	7,7	8,4	11,3	11,6	11,0	5,9	6,2	5,6
Estonia	13,6	14,5	12,7	7,9	8,8	7,1	12,5	13,1	11,8
Latvia	13,7	14,4	12,9	8,9	9,1	8,7	15,4	17,6	13,1
Lithuania	16,4	18,6	14,1	8,3	8,2	8,3	15,4	17,8	13,0
Poland	16,1	14,4	18,2	17,8	16,6	19,2	9,7	9,0	10,5
Finland	9,8	9,1	10,6	8,4	8,2	8,6	7,8	8,4	7,1
Sweden	5,6	5,9	5,3	7,7	7,7	7,6	7,5	7,6	7,5

Szczegółowe analizy zharmonizowanej stopy bezrobocia wymagają dodatkowo uwzględnienia wieku osoby pozostającej bez pracy, przynajmniej w tym najprostszym podziale, na dwie grupy wiekowe – do 25 lat i powyżej 25 lat. Jest to dosyć istotny podział, gdyż uwzględnia społeczne uwarunkowania rynku pracy, niechęć pracodawców do zatrudniania młodych pracowników i wyraźnie zawężony sektor miejsc pracy dostępnych dla absolwentów szkół średnich i wyższych. Jest to zapewne związane również z mobilnością lub jej brakiem w tych grupach wiekowych, co przekłada się na strukturę rynku pracy.

W tab. 2 zawarto współczynniki poziomu bezrobocia z uwzględnieniem podziału wiekowego na dwie wspomniane grupy. Wyraźnie widoczna jest różnica w ich wartościach. Dla młodych osób poziom bezrobocia jest ponad dwukrotnie wyższy w skali całej UE, natomiast w poszczególnych państwach różni się zdecydowanie. W Danii i w Niemczech dysproporcja jest najniższa (jak również poziom bezrobocia jest najniższy). W państwach dawnej wspólnoty rosyjskiej i w Polsce różnice są najbardziej widoczne, również wartości współczynników są najwyższe i kilkakrotnie przekraczają poziom w państwach starej unii.

Tabela 2

Stopa bezrobocia w wybranych grupach wieku w państwach nadbałtyckich
w latach 2000, 2005, 2011

	2000 r.		2005 r.		2011 r.	
	do 25 lat	25 do 74 lat	do 25 lat	25 do 74 lat	do 25 lat	25 do 74 lat
EU 27	17,5	7,5	18,8	7,7	21,4	8,3
Denmark	6,2	4	8,6	4,2	14,2	6,3
Germany	8,7	7,9	15,6	10,7	8,6	5,6
Estonia	24,4	12,3	15,9	7	22,3	11,3
Latvia	21,4	12,7	13,6	8,3	29,1	13,8
Lithuania	30,6	14,6	15,7	7,6	32,9	13,8
Poland	35,1	13,3	36,9	15,1	25,8	8
Finland	21,4	8,1	20,1	6,8	20,1	6,1
Sweden	10,5	5	22,6	5,7	22,9	5,2

Widoczne jest, że spowodowany kryzysem wzrost bezrobocia najmocniej dotknął młodych obywateli Unii Europejskiej. Natężenie bezrobocia wśród Europejczyków w wieku 15-24 lata jest nie tylko znacznie wyższe niż w przypadku ogółu mieszkańców Wspólnoty, ale również wzrasta dużo szybciej, np. w I kwartale 2009 r. niemal co piąty młody obywatel UE nie miał pracy. Stopa bezrobocia wśród osób w wieku 15-24 lata ukształtowała się na poziomie 18,9%, aby w 2011 r. osiągnąć poziom 21,4% – to o 6,5 punktu procentowego więcej niż w I kwartale 2008 r. W tym samym okresie natężenie bezrobocia wśród ogółu mieszkańców UE zwiększyło się o 2,6 punktu procentowego i osiągnęło poziom 9,7% – ponad dwa razy mniej niż w przypadku najmłodszych uczestników rynku pracy.

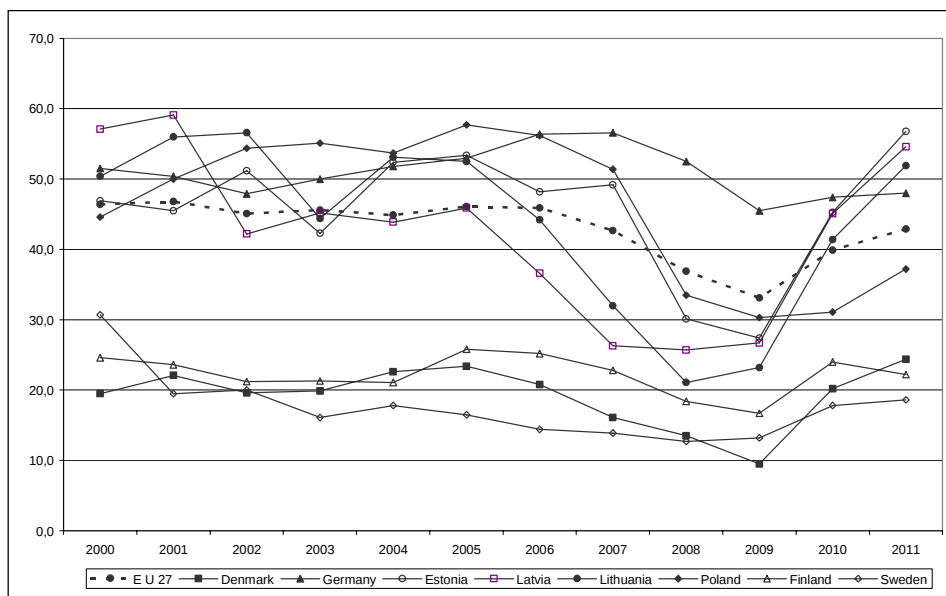
Jak wspomniano wyżej, w analizowanym okresie odsetek młodych Europejczyków poszukujących zatrudnienia wzrósł we wszystkich państwach członkowskich Wspólnoty. Najbardziej ich grono zwiększyło się w republikach bałtyckich. Stopa bezrobocia Łotyszy w wieku 15-24 lata wzrosła o 9,7 punktów procentowych, a wśród Estończyków i Litwinów wskaźnik ten zwiększył się o 8,7 punktów procentowych. Kraje bałtyckie to jednocześnie państwa, w których w 2011 r. natężenie bezrobocia wśród najmłodszych uczestników rynku pracy osiągnęło jeden z najwyższych poziomów w Unii Europejskiej – na Łotwie 29,1,2%, Litwie 32,9%, w Estonii 22,3%. Przekraczająca 20 punktów procentowych stopa bezrobocia wśród młodych obywateli dotyczy również Szwecji (21,9%) i Polski (25,2%). Analizując zależność poziomu bezrobocia od wieku potencjalnego pracownika warto wspomnieć, że grupą najbardziej narażoną na wzrost skali bezrobocia są osoby w wieku powyżej 50 lat. O tym, że stopa bez-

robocia w grupie wieku 50+ różni się od poziomu bezrobocia ogółem w całej zbiorowości pracowników można się przekonać porównując wartości obydwu wskaźników w dowolnie wybranym czasie. Szczegółowe analizy prowadzone w ramach szerszych badań problemu bezrobocia [2] pozwalają na stwierdzenie, że poziom bezrobocia dla ogółu pracujących jest wyższy niż w grupie wiekowej 50-64 lata. 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. Jak jednak wskazuje praktyka, jest to działanie pozorne, gdyż niekonięcznie 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 wieku 50+ maleje, a w przypadku ogółu pracowników pozostaje bez zmian lub wręcz rośnie.

2. Bezrobocie długoterminowe

Jak wspomniano wcześniej, odrębnym problemem w zagadnieniach dotyczących sfery bezrobocia jest bezrobocie długoterminowe, czyli sytuacja, gdy pracownik poszukuje bezskutecznie pracy przez okres dłuższy niż rok. Bezrobocie długookresowe jest szczególną formą bezrobocia i różni się w znacznym stopniu od pozostałych jego form. Można to przedstawić następująco: po pierwsze, następuje swoista „profesjonalizacja” statusu bezrobotnego, czyli bezrobocie staje się w coraz większym stopniu sposobem na życie, po wtóre, aktywizacja bezrobotnych długotrwale jest trudniejsza niż bezrobotnych przejściowo. Długotrwale bezrobocie w dużym stopniu jest zdeterminowane przez płeć. Znaczną część z tej grupy bezrobotnych stanowią kobiety, ponieważ są uważane przez pracodawców za pracowników mniej dyspozycyjnych i bardziej kłopotliwych (urlopy macierzyńskie, wychowawcze, zwolnienia na opiekę nad chorym dzieckiem) oraz zakres dyspozycyjności zawodowej jest w przypadku kobiet węższy (pracują w mniejszej liczbie zawodów, zwłaszcza związanych z przemysłem ciężkim). Wiek i staż pracy, jako silne cechy ze sobą skorelowane, wywierają podobny wpływ na czas pozostawania bez pracy. Wśród ludzi bezrobotnych dużą grupą są ludzie młodzi niemający więcej niż 25 lat oraz osoby, które przekroczyły 50. rok życia. Zdecydowana większość długotrwale bezrobotnych to osoby, które wcześniej pracowały. W przeważającej części utracili oni pracę w związku z likwidacją zakładu lub stanowiska pracy. Najczęściej byli pracownikami przemysłu i budownictwa, rzadziej handlu prywatnego. Długoterminowość bezrobocia wynika z nieodpowiednich kwalifikacji lub wręcz ich braku bądź z niedopasowania kwalifikacji do wymogów rynku pracy. Bezrobocie dłu-

goterminowe jest problemem nie tylko polskiego rynku pracy, lecz dotyka również rynek pracy w zasadzie wszystkich państw europejskich. W przypadku państw nadbałtyckich tempo zmian stopy bezrobocia długoterminowego dla ogółu zasobów pracy przedstawia rys. 4.



Rys. 4. Stopa bezrobocia długoterminowego ogółem dla grupy państw nadbałtyckich w latach 2000-2011

Szczegółowa analiza linii trendu stopy bezrobocia długoterminowego pozwala na sformułowanie ciekawych wniosków:

- Wyraźnie niższy jest poziom bezrobocia długoterminowego w państwach skandynawskich. Związane to jest m.in. z lepiej prowadzonymi działaniami mającymi na celu ograniczanie czasu pozostawania bez pracy osób zdolnych ją wykonywać.
- Dynamika bezrobocia długoterminowego dla nadbałtyckich państw skandynawskich charakteryzuje się najmniejszą zmiennością i odsetek osób pozostających bez pracy dłużej niż rok wśród ogółu bezrobotnych oscyluje wokół 20%.
- Wśród państw „starej Unii” zdecydowanie najwyższy jest odsetek bezrobotnych długookresowo w Niemczech. Przekracza on dwukrotnie poziom współczynnika dla państw skandynawskich. Powodem takiej sytuacji jest zapewne „nadopiekuńczość” państwa i szeroko rozwinięty system zabezpieczenia społecznego widoczny w zasadach przyznawania zasiłków socjalnych dla bezrobotnych, gdy zasiłek jest na tyle wysoki, że wystarcza na zaspoko-

jenie podstawowych potrzeb i wówczas status bezrobotnego staje się wygodnym sposobem na życie. Właśnie w odniesieniu do Niemiec można mówić o swoistej „profesjonalizacji” bezrobocia.

- W grupie państw nadbałtyckich nowo wstępujących do UE wysoki poziom bezrobocia długoterminowego nie jest niespodzianką. Przyczyniają się do tego przemiany restrukturyzacyjne w gospodarkach tych państw i obecnie nie ma skutecznego sposobu ograniczania bezrobocia, zwłaszcza długoterminowego.

Podsumowanie

Zaprezentowane 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.

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 konsekwencje zmiennej dynamiki parametrów rynku pracy spowodują zanik płynności finansowej gospodarstw domowych, co wywoła problemy ze ściągalskością kredytów krótko- i długoterminowych, a zwłaszcza kredytów hipotecznych. Ponadto 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. Problemy rynku pracy w przypadku ich nieefektywnego rozwiązania będą zatem obejmować coraz szersze dziedziny gospodarki.

Literatura

1. *Aktywność ekonomiczna ludności Polski*, GUS, Warszawa 2009.
2. Balcerowicz-Szkutnik M., *Bezrobocie i jego ekonomiczne konsekwencje – analiza porównawcza dla wybranej grupy państw UE*, Studia Ekonomiczne, nr 78, Uniwersytet Ekonomiczny, Katowice 2011.
3. *Popyt na pracę w 2008*, GUS, Warszawa 2009.

4. *Przejsście z pracy na emeryturę*, GUS, Warszawa 2007.
5. *Unia Europejska – wskaźniki krótkookresowe*, GUS, Warszawa 2007.
6. www.diagnoza.com/pliki/raporty/diagnoza_raport_2009.pdf
7. www.rynekpracy.pl
8. www.stat.gov.pl/eutostat

DETERMINANTS OF UNEMPLOYMENT SELECTED EU COUNTRIES – – STATISTICAL ANALYSIS

Summary

In the article there are presented findings of detailed analyses of dynamics one of basic parameters which characterize labour market, that is unemployment rate and long-term unemployment rate for the group new EU members. Analysis includes construction of trend function of aforesaid parameters and prognoses of their values for forthcoming years which are assigned by a several independent methods. Separate part of elaboration is estimation of influence parameters describing the labour market on basic economic and social processes

Anna Sączewska-Piotrowska

Uniwersytet Ekonomiczny w Katowicach

PROGNOSTYCZNY WARIANT UBÓSTWA DLA GOSPODARSTW DOMOWYCH MAKROREGIONU POŁUDNIOWEGO

Wprowadzenie

Analiza sfery ubóstwa jest najczęściej przeprowadzana w celu umożliwienia porównań oraz wskazania grup jednostek najbardziej zagrożonych biedą. Celem niniejszego opracowania jest ocena sfery ubóstwa w makroregionie południowym* oraz sformułowanie wniosków dotyczących kształtowania się badanego zjawiska w najbliższych latach. W pracy dokonano szerokiej analizy, obejmującej z jednej strony zjawisko ubóstwa, z drugiej strony uwzględniającej sytuację przeciwną do biedy, mianowicie dobrobyt. Badaniem objęto również nierówności dochodowe będące *de facto* podstawą określania ubóstwa. Wartości wybranych mierników wyznaczone dla regionu południowego porównano z wartościami mierników uzyskanymi dla całej Polski. Dzięki temu uzyskano odpowiedź na pytanie, czy region południowy jest zagrożony ubóstwem w mniejszym czy większym stopniu niż reszta kraju i czy tym samym powinien być adresatem szczególnych działań państwa w walce z ubóstwem.

1. Przyjęte założenia metodologiczne

Badanie nierówności dochodowych i ubóstwa gospodarstw domowych makroregionu południowego przeprowadzono wykorzystując dane za lata 2000-2009 zawarte w bazie projektu „Diagnoza społeczna”. Podstawowe informacje dotyczące badanych gospodarstw domowych w latach 2000-2009 przedstawiono w tab. 1.

* W niniejszej pracy pojęcia „region” i „makroregion” są używane zamiennie.

Tabela 1

Liczba gospodarstw domowych objętych badaniem w regionie południowym
w latach 2000-2009

Wyszczególnienie	Liczba gospodarstw				
	2000 r.	2003 r.	2005 r.	2007 r.	2009 r.
Ogółem	344	562	513	871	2110
Liczba osób w gospodarstwach	1181	1915	1683	2734	6232
Przeciętna liczba osób w gospodarstwie	3,43	3,41	3,28	3,14	2,95

Źródło: Opracowanie własne na podstawie: [5].

W latach 2000-2009 grupa badanych gospodarstw zwiększyła się – wyjątkiem był 2005 r., w którym zanotowano spadek liczebności grupy. Można założyć, że w tych latach następował stopniowy spadek średniej liczby osób w gospodarstwie (od 3,43 w 2000 r. do 2,95 w 2009 r.).

Analiza nierówności dochodowych i ubóstwa wymaga dokonania wielu założeń metodologicznych. Przyjęto, że miernikiem zamożności są dochody netto uzyskiwane przez gospodarstwa domowe w regionie południowym w lutym 2000 r., 2003 r., 2005 r., 2007 r. i 2009 r. Dla zachowania porównywalności sytuacji gospodarstw o różnym składzie demograficznym obliczono dochody ekwiwalentne stosując zmodyfikowaną skalę ekwiwalentności OECD typu 0,5/0,3. Przyjęto również, że dochody przeliczone na jednostkę ekwiwalentną są ważone liczbą osób w gospodarstwie domowym.

Charakterystykę rozkładów dochodów ekwiwalentnych przeprowadzono wykorzystując podstawowe miary opisowe oraz wybrane miary nierówności i dobrobytu, takie jak: współczynniki Giniego, Schutza, Atkinsona oraz indeks Sena. W tab. 2 przedstawiono postaci wzorów wykorzystane w przeprowadzonej analizie.

Tabela 2

Miary nierówności rozkładów dochodów oraz miary dobrobytu

Nazwa	Postać miernika	Oznaczenia
1	2	3
Współczynnik Giniego*	$G = 1 + \frac{1}{n} - \frac{2}{n^2 \bar{y}} \sum_{i=1}^n (n+1-i)y_i,$	n – liczba gospodarstw domowych $\bar{y}$ – średnia arytmetyczna w rozkładzie dochodów y_i – dochód i-tej osoby

cd. tabeli 2

1	2	3
Współczynnik Schutza	$S = \frac{1}{2\bar{y}} \frac{\sum_{i=1}^n y_i - \bar{y} }{n}$	Oznaczenia jw.
Indeks Atkinsona**	$IA = 1 - \frac{y_g}{\bar{y}}$	y_g – średnia geometryczna w rozkładzie dochodów $\bar{y}$ – jw.
Indeks Sena	$IS = \bar{y}(1 - G),$	Oznaczenia jw.

* Wzór dla niegrupowanych jednostek, uporządkowanych według niemalejących dochodów.

** Wzór przyjmuje taką postać przy założeniu, że funkcja użyteczności dochodów poszczególnych jednostek jest typu Bernoulliego, tj. $u(y) = u_1(y) = a + b \ln y$.

Źródło: Opracowanie własne na podstawie: [1; 3; 7].

Analizując ubóstwo obiektywne przyjęto taką samą skalę ekwiwalentności, jaką zaproponowano przy badaniu nierówności dochodowych. W badaniu zastosowano wykorzystywaną przez Eurostat relatywną linię ubóstwa obliczaną jako 60% mediany dochodów ekwiwalentnych ogółu gospodarstw. Dysponując wyznaczonymi granicami ubóstwa obliczono wskaźniki pozwalające na ocenę zasięgu, głębokości i dotkliwości ubóstwa w danym roku badania. Analizy dokonano za pomocą klasy mierników ubóstwa zaproponowanej przez Fostera, Greera i Thorbecke. Grupę mierników FGT można zdefiniować następująco [2]:

$$P_{\alpha}(y, y^*) = FGT(\alpha) = \frac{1}{n} \sum_{i=1}^q \left(\frac{y^* - y_i}{y^*} \right)^{\alpha} \quad (1)$$

gdzie:

y – wektor dochodów,

y^* – granica ubóstwa,

α ($\alpha \geq 0$) – parametr miernika,

n – liczba gospodarstw domowych,

q – liczba ubogich gospodarstw domowych.

Miernik P_{α} , zaproponowany przez Fostera, Greera i Thorbecke, w zależności od przyjętej wartości parametru α można różnie zinterpretować. Przyjmując w formule (1) $\alpha = 0$ otrzymujemy wskaźnik równy stopie ubóstwa, który informuje o zasięgu ubóstwa (o udziale ubogich gospodarstw domowych w danej populacji), przyjmując $\alpha = 1$ otrzymujemy miernik głębokości ubóstwa, natomiast podstawiając $\alpha = 2$ uzyskujemy wskaźnik dotkliwości ubóstwa.

Głębokość ubóstwa jest również często mierzona wskaźnikiem luki dochodowej (luki ubóstwa, średniej luki dochodowej). Wskaźnik ten nie należy do klasy miar zaproponowanych przez Fostera, Greera i Thorbecke, a od wskaźnika głębokości ubóstwa $FGT(1)$ różni się tym, że mówi o przeciętnym zubożeniu w grupie biednych, a nie w całej populacji. Formuła pozwalająca na wyznaczenie miernika luki dochodowej przyjmuje postać:

$$PG = \frac{1}{q} \sum_{i=1}^q \left(\frac{y^* - y_i}{y^*} \right) = \frac{y^* - \bar{y}_q}{y^*} \quad (2)$$

gdzie:

$\bar{y}_q$ – średni dochód w grupie ubogich gospodarstw domowych.

Ze względu na ograniczenia objętościowe opracowania przedstawione miary nierówności, dobrobytu i ubóstwa nie zostaną bliżej scharakteryzowane. Szczegółowe informacje dotyczące zastosowanych miar można znaleźć m.in. w pracach [4; 6].

2. Nierówności dochodowe i dobrobyt ekonomiczny gospodarstw domowych

W celu przeprowadzenia analizy rozkładów dochodów ekwiwalentnych w latach 2000-2009 wyznaczono podstawowe parametry opisowe charakteryzujące te rozkłady (tab. 3).

Tabela 3

Parametry opisowe rozkładów dochodów ekwiwalentnych w regionie południowym w latach 2000-2009

Wyszczególnienie	Lata				
	2000	2003	2005	2007	2009
Średnia arytmetyczna	765,02	877,73	1061,88	1209,25	1463,61
Mediana	692,31	800,00	928,57	1050,00	1317,07
Współczynnik zmienności (%)	49,81	51,71	68,94	60,91	58,59
Współczynnik asymetrii	1,747	1,504	6,000	2,709	2,351

Źródło: [5].

W latach 2000-2009 wartości średniej arytmetycznej i mediany dochodów ekwiwalentnych rosły z okresu na okres. W badanym okresie nastąpił wzrost przeciętnych dochodów o 17,61 %, natomiast mediany dochodów o 17,44%*. Najwyższy wzrost wartości średniej i mediany dochodów ekwiwalentnych w stosunku do poprzedniego okresu nastąpił w 2009 r. i wyniósł odpowiednio 21,03% i 25,44%. W badanym okresie współczynnik zmienności dochodów ekwiwalentnych utrzymywał się na dosyć wysokim poziomie – od blisko 50% do prawie 69%. Rozkłady dochodów we wszystkich analizowanych latach cechowały się asymetrią prawostronną, co oznacza, że większość członków gospodarstw osiągała dochody poniżej średniej arytmetycznej wyznaczonej dla danego rozkładu.

W tab. 4 przedstawiono wybrane miary nierówności rozkładów dochodów ekwiwalentnych oraz miary dobrobytu, które wyznaczono dla pięciu etapów badania w latach 2000-2009.

Tabela 4

Miary nierówności rozkładów dochodów ekwiwalentnych oraz miary dobrobytu
w latach 2000-2009 w regionie południowym

Wyszczególnienie	Lata				
	2000	2003	2005	2007	2009
Współczynnik Giniego	0,257	0,272	0,292	0,297	0,290
Współczynnik Schutza	0,179	0,193	0,202	0,210	0,203
Współczynnik Atkinsona	0,107	0,122	0,141	0,141	0,137
Indeks Sena	568,73	638,57	751,90	850,66	1039,81

Źródło: [5].

W latach 2000-2009 wartości obliczonych współczynników Giniego, Schutza oraz Atkinsona wzrosły, co wskazuje na wzrost nierówności dochodowych w badanym okresie. Wyjątkiem okazał się 2009 r., w którym zanotowano spadek wartości współczynników w porównaniu do poprzedniego okresu.

W badanym okresie współczynnik Giniego utrzymywał się na umiarkowanym poziomie, przyjmując największą wartość równą 0,297 w 2007 r. i wskazując tym samym na największe nierówności dochodowe. Uzyskana w 2007 r. wartość współczynnika informuje, że przeciętna absolutna różnica pomiędzy dochodami losowo wybranej pary osób stanowiła 59,4% dochodu średniego (podwojona wartość współczynnika Giniego).

* Analizując dynamikę dochodów w latach 2000-2009 brano pod uwagę dochody w ujęciu nominalnym.

Współczynnik Schutza przyjął w 2007 r. najwyższą wartość równą 0,210, co oznacza, że w tym roku występowały największe nierówności dochodowe. Współczynnik ten interpretuje się w następujący sposób: jeśli całą populację członków gospodarstw podzielimy na dwie grupy: grupę, składającą się z osób o dochodach ekwiwalentnych poniżej lub równych średniej, drugą, składającą się z osób o dochodach powyżej średniej, to 21% ogólnego dochodu powinno być transferowane z grupy zamożniejszej do biedniejszej, aby obie grupy miały dokładnie taki sam dochód przeciętny, tzn. aby zniknęły nierówności dochodowe.

Na największe nierówności dochodowe w 2005 r. i 2007 r. wskazuje współczynnik Atkinsona, który w obydwu okresach przyjął wartość równą 0,141. Uzyskany wynik oznacza, że poświęcenie przez każdą osobę kwoty rzędu 14,1% dochodu przeciętnego (około 150 zł w 2005 r. i około 170,5 zł w 2007 r.) zlikwidowałoby całkowicie nierówności, bez zmniejszania dobrobytu społecznego.

W badanych latach wartość indeksu dobrobytu Sena rosła z okresu na okres. W latach 2000-2009 nastąpił wzrost indeksu o 16,28%, przy czym najwyższym wzrostem wartości indeksu w stosunku do poprzedniego okresu cechował się 2009 r. (wzrost o 22,2%), natomiast najniższym wzrostem 2003 r. (wzrost o 12,3%).

3. Ocena poziomu ubóstwa obiektywnego

Sferę ubóstwa w makroregionie południowym przeanalizowano uwzględniając granicę ubóstwa równą 60% mediany dochodów ekwiwalentnych ogółu gospodarstw w Polsce. Wartości linii ubóstwa w poszczególnych latach w Polsce przedstawiono w tab. 5. Linie te stanowiły podstawę obliczenia wskaźników ubóstwa w regionie południowym (tab. 6).

Tabela 5

Wartości granic ubóstwa obiektywnego w Polsce
w latach 2000-2009 (w zł)

Lata	Granica ubóstwa
2000	400,00
2003	454,55
2005	530,77
2007	600,00
2009	738,46

Źródło: [5].

Tabela 6

Wskaźniki ubóstwa obiektywnego w regionie południowym
w latach 2000-2009 (% osób)

Wyszczególnienie	Lata				
	2000	2003	2005	2007	2009
Stopa ubóstwa	11,85	14,26	15,21	13,53	14,81
Wskaźnik luki dochodowej	26,09	28,79	26,78	29,48	27,43
Wskaźnik głębokości ubóstwa	3,09	4,10	4,07	3,99	4,06
Wskaźnik dotkliwości ubóstwa	1,17	1,72	1,77	1,65	1,70

Źródło: [5].

Analizując mierniki ubóstwa można zauważyć, że przyjmowały one najniższe wartości w 2000 r. Wyznaczona dla 2000 r. stopa ubóstwa informuje, że odsetek osób w gospodarstwach domowych, których poziom dochodów był niższy od obiektywnej granicy ubóstwa wyniósł 11,85%. Obliczona wartość wskaźnika luki dochodowej informuje, że przeciętny dochód ekwiwalentny osób w gospodarstwach ubogich jest o 26,09% niższy od granicy ubóstwa. Indeks głębokości ubóstwa informuje natomiast, że przeciętny dochód ekwiwalentny osób we wszystkich gospodarstwach domowych jest o 3,09% niższy od linii ubóstwa. Wskaźnik dotkliwości ubóstwa, osiągający najmniejszą wartość w 2000 r. informuje, że w tym roku dochody ekwiwalentne osób w gospodarstwach ubogich są najmniej oddalone od wyznaczonej granicy ubóstwa.

Dwa z obliczonych wskaźników ubóstwa przyjmowały najwyższe wartości w 2005 r. – stopa ubóstwa i dotkliwość ubóstwa, natomiast głębokość ubóstwa była największa w 2003 r., a luka dochodowa w 2007 r. Na podstawie uzyskanych wyników nie można więc jednoznacznie wskazać roku, w którym sytuacja materialna gospodarstw domowych w regionie południowym była najgorsza.

W tab. 7 przedstawiono wybrane charakterystyki rozkładów dochodów osób w gospodarstwach domowych zagrożonych ubóstwem.

Tabela 7

Wybrane charakterystyki rozkładów dochodów ekwiwalentnych osób należących do sfery ubóstwa obiektywnego w regionie południowym w latach 2000-2009

Miary	Lata				
	2000	2003	2005	2007	2009
Średnia arytmetyczna	295,64	323,67	388,64	423,12	535,90
Współczynnik zmienności (%)	23,80	27,35	28,89	26,54	27,42
Współczynnik Giniego	0,133	0,155	0,162	0,148	0,151
Indeks Sena	256,42	273,65	325,78	360,63	454,73

Źródło: [5].

Można zauważyć, że oprócz niskich średnich dochodów ekwiwalentnych, rozkład dochodów osób w gospodarstwach ubogich charakteryzuje się małymi wartościami współczynników zmienności i Giniego oraz indeksu Sena. Wartości wspomnianych mierników są w każdym z badanych okresów około dwukrotnie niższe w porównaniu z wartościami obliczonymi dla wszystkich badanych osób w gospodarstwach w regionie południowym. Mniejsze wartości miar są oczywistą konsekwencją braku wysokich dochodów wśród ubogich – dochody są od góry ograniczone wartością równą linii ubóstwa.

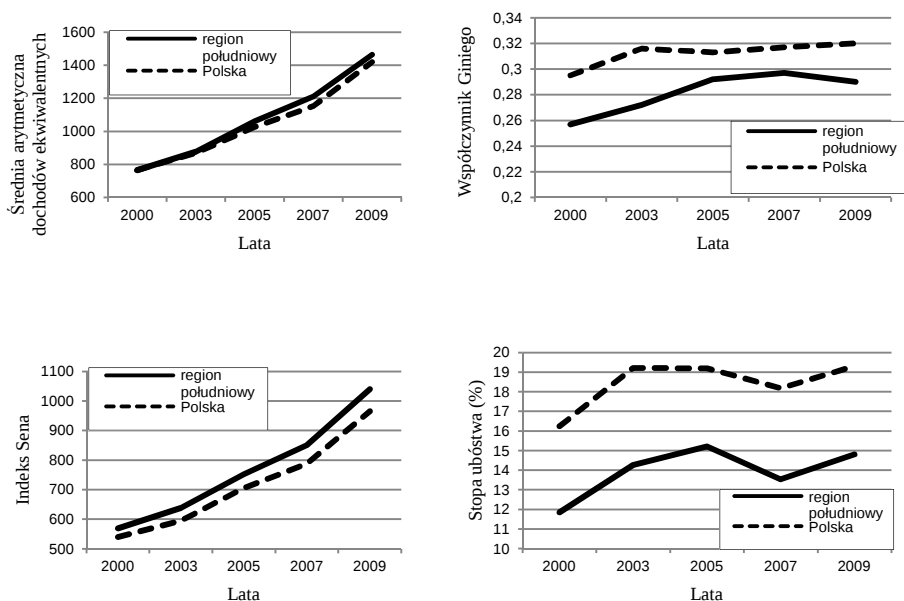
W latach 2000-2009 przeciętne dochody ekwiwalentne osób zaliczanych do sfery ubóstwa rosły z okresu na okres średnio o 16,03%, przy czym największy wzrost (o 26,65%) w porównaniu z poprzednim okresem zanotowano w 2009 r. Najmniejszy wzrost przeciętnych dochodów wystąpił w latach 2003 i 2007 (w porównaniu do poprzednich okresów), wynosząc w obydwu przypadkach około 9%.

W latach 2000-2005 wartości współczynników zmienności i Giniego rosły z okresu na okres. W 2007 r. nastąpił spadek wartości obydwu rozważanych współczynników, natomiast w 2009 r. miał miejsce ponowny wzrost ich wartości.

W badanym okresie indeks dobrobytu Sena cechował się tendencją rosnącą – zanotowano wzrost indeksu z okresu na okres o 15,4%. Najwyższy wzrost omawianej miary w stosunku do poprzedniego okresu wystąpił w 2009 r. (wzrost o 26,09%), natomiast najmniejszy wzrost w porównaniu do poprzedniego okresu w 2003 r. (wzrost o 6,72%).

4. Dobrobyt oraz ubóstwo ekonomiczne w Polsce i w makroregionie południowym

Uzyskane wartości mierników nierówności dochodowych, dobrobytu oraz ubóstwa w makroregionie południowym porównano z wartościami mierników uzyskanymi dla całej Polski (rys. 1).

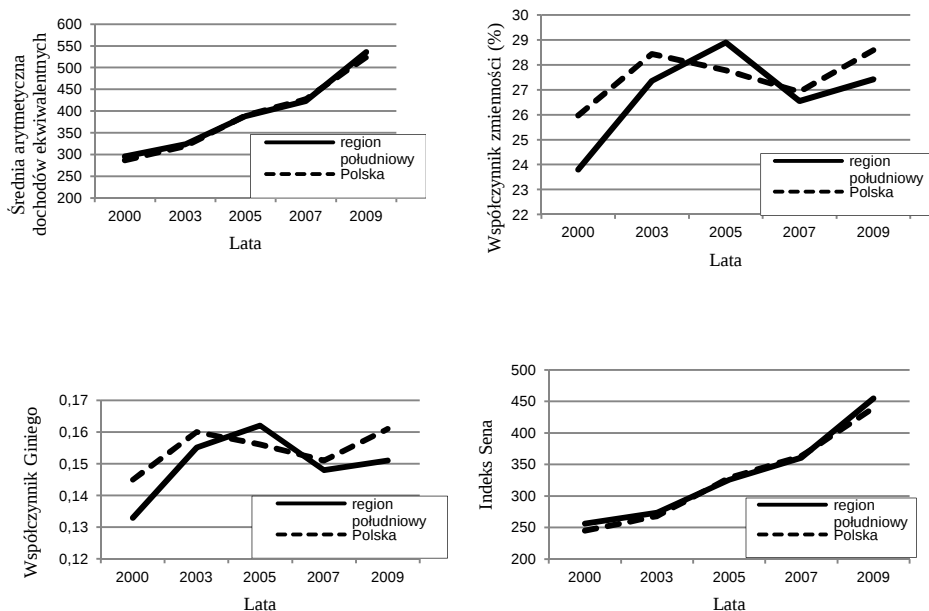


Rys. 1. Porównanie wybranych miar nierówności dochodowych, dobrobytu i ubóstwa w Polsce oraz w regionie południowym w latach 2000-2009

Źródło: [5].

W latach 2000-2009 zarówno miary nierówności dochodowych (współczynnik Giniego), dobrobytu (indeks Sena), jak i zasięgu ubóstwa (stopa ubóstwa) przyjmowały bardziej korzystne wartości w regionie południowym niż w Polsce. Również podstawowa miara – średnie dochody ekwiwalentne była wyższa w regionie południowym we wszystkich latach badania.

Wybrane charakterystyki rozkładów dochodów ekwiwalentnych osób należących do sfery ubóstwa w regionie południowym porównano z charakterystykami wyznaczonymi dla całego kraju (rys. 2).



Rys. 2. Zagrożeni ubóstwem – porównanie wybranych miar nierówności dochodowych i dobrobytu w Polsce oraz w regionie południowym w latach 2000-2009

Źródło: [5].

W przypadku gospodarstw ubogich nie można wysnuć tak jednoznacznego wniosku, jak w przypadku ogółu gospodarstw. Średnia arytmetyczna dochodów osób ubogich oraz indeks Sena w Polsce i w regionie południowym kształtowały się w badanych latach na zbliżonym poziomie. W przypadku współczynników zmienności i Giniego w większości badanych okresów korzystniejsze wartości występowały w regionie południowym. Jedynie w 2005 r. wartości wspomnianych współczynników były większe dla makroregionu południowego niż dla Polski.

Podsumowanie

Podsumowując wyniki analizy sfery ubóstwa w makroregionie południowym można stwierdzić, że region ten jest w lepszej sytuacji materialnej w porównaniu do Polski. Z przeprowadzonego badania wynika, że w latach 2000-

-2009 różnica pomiędzy wartościami mierników nierówności, dobrobytu i ubóstwa dla Polski oraz dla regionu południowego utrzymywała się na stałym poziomie, co pozwala wysnuć wniosek, że w najbliższych latach sytuacja nie powinna ulec zmianie. Niezbędnym warunkiem jest brak wystąpienia niespodziewanych zdarzeń w sferze gospodarczej, społecznej itd., które by mogły zakłócić tę stabilną sytuację.

W badanych latach rosnące wartości indeksu dobrobytu Sena osób ubogich w regionie południowym utrzymywały się na poziomie zbliżonym do poziomu całego kraju. Można się więc również spodziewać, że tendencja ta (przy niezmienionych warunkach) w najbliższym czasie nie ulegnie zmianie. Trudne do przewidzenia jest natomiast kształtowanie się wartości współczynnika zmienności oraz współczynnika Giniego wśród ubogich. Wartości tych mierników utrzymywały się w regionie południowym w większości badanych lat poniżej wartości obliczonych dla Polski, jedynie w 2005 r. miała miejsce odwrotna sytuacja. Traktując 2005 r. jako rok wyjątkowy, na podstawie uzyskanych wyników można się spodziewać, że region południowy będzie się charakteryzował mniejszym zróżnicowaniem dochodów niż Polska.

Należy zaznaczyć, że duży wpływ na uzyskane wyniki dotyczące sfery ubóstwa miało przyjęcie granicy ubóstwa wyznaczonej dla całego kraju. Zabieg ten był celowy – umożliwił dokonanie porównań z całym krajem. Z drugiej strony można się spodziewać, że granica ubóstwa liczona tylko na podstawie dochodów ekwiwalentnych w regionie południowym przyjęłaby wyższą wartość (mediana dla regionu była wyższa niż wyznaczona dla kraju), co skutkowałoby mniejszą stopą ubóstwa, lecz wyższymi wartościami mierników głębokości ubóstwa.

Literatura

1. Atkinson A.B., *On the measurement of inequality*, „Journal of Economic Theory” 1970, Vol. 2.
2. Foster J., Greer J., Thorbecke E., *A class od decomposable poverty measures*, „Econometrica” 1984, Vol. 52.
3. Kondor Y., *An old-new measure of income inequality*, „Econometrica” 1971, Vol. 39.
4. Kot S.M., *Ekonometryczne modele dobrobytu*, PWN, Warszawa-Kraków 2000.
5. Rada Monitoringu Społecznego: *Diagnoza społeczna 2009: zintegrowana baza danych*, <http://www.diagnoza.com> (05.01.2010).

6. Rusnak Z., *Statystyczna analiza dobrobytu ekonomicznego gospodarstw domowych*, Akademia Ekonomiczna, Wrocław 2007.
7. Sen A., *Poverty: an ordinal approach to measurement*, „Econometrica” 1976, Vol. 44.

PROGNOSTIC VARIANT OF POVERTY FOR HOUSEHOLDS OF SOUTHERN MACROREGION

Summary

The aim of this study is evaluation of poverty range in southern macroregion and formulate conclusions for the form of the studied phenomenon in the coming years. The paper presents a deep analysis, involving on the one hand poverty itself, on the other hand taking into the consideration the situation opposite to poverty – welfare. The income inequality, which is *de facto* the basis of defining poverty, was also taken into the consideration. The values of selected indicators calculated for southern macroregion was also compared with the values of indicators calculated for Poland.

Mirosław Wójciak

Uniwersytet Ekonomiczny w Katowicach

METODY BUDOWY DŁUGOTERMINOWYCH PROGNOZ PRZEDZIAŁOWYCH ROZWOJU NOWYCH ZJAWISK

Wprowadzenie

Budowa prognoz rozwoju nowych zjawisk, np. nowych technologii, nowych produktów jest trudna ze względu na brak danych historycznych oraz długoterminowy charakter prognozy. W przypadku, gdy nie dysponuje się danymi dotyczącymi kształtowania się danego zjawiska w przeszłości można zbudować subiektywne modele tendencji rozwojowej, których parametry są szacowane na podstawie opinii ekspertów [1; 2; 6]. W przypadku dalekiego horyzontu prognozy nie jest spełnione jedno z klasycznych założeń teorii predykcji, dotyczące stabilności lub prawie stabilności prawidłowości ekonomicznej w czasie [5]. Budując długoterminową prognozę należy uwzględnić zmienność otoczenia (np. sytuacji ekonomicznej, politycznej, prawnej, społecznej, technologicznej oraz stanu naturalnego środowiska). Problem ten może zostać częściowo rozwiązany poprzez budowę prognoz wariantowych, uwzględniających przyjęte scenariusze rozwoju otoczenia, jednak w tym przypadku otrzymuje się tylko prognozy punktowe, które do konstrukcji scenariuszy rozwoju mogą być niewystarczające. Analizę należałoby wzbogacić o przedział prognoz, który by uwzględniał zmiany w rozwoju rozpatrywanego zjawiska ze względu na zmianę poziomów czynników kluczowych. Rozpiętość zbudowanego przedziału prognoz można w tym przypadku traktować jako stopień niepewności prognozy. Czynniki kluczowe można zdefiniować jako zmienne opisujące powolne, regularne zmiany otoczenia, które mają istotny wpływ na rozwój analizowanego zjawiska. Na podstawie danych pozyskanych od ekspertów, dotyczących m.in. prawdopodobieństwa zajścia danego poziomu czynnika, wpływu poziomu danego czynnika na rozpatrywane zjawisko, możliwa jest budowa prognoz przedziałowych z określonym prawdopodobieństwem.

Celem opracowania jest przedstawienie sposobu budowy długoterminowych prognoz przedziałowych, uwzględniających zmiany otoczenia analizowanego zjawiska. Zaprezentowane metody mogą mieć zastosowanie w foresightach, których nadrzędnym celem jest konstrukcja scenariuszy rozwoju sytuacji w stosunkowo dalekiej perspektywie (zwykle 20-30 lat). W opracowaniu porównano przedziały prognoz uzyskane za pomocą: standardowej niepewności prognoz, agregacji czynników kluczowych i analizy symulacyjnej. Przykłady zastosowania omawianych metod przedstawiono na podstawie danych uzyskanych w ramach projektu „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050” realizowanego w Głównym Instytucie Górnictwa*.

1. Subiektywne modele tendencji rozwojowej

Aby uzyskać prognozy przedziałowe, w pierwszej kolejności należy zbudować prognozy punktowe o zadanym horyzoncie, które w dalszej części opracowania będą nazywane prognozami bazowymi. W tym celu można ekstrapolować oszacowany model tendencji rozwojowej na podstawie danych historycznych. W przypadku, gdy prognozuje się zjawisko nowe, dla którego nie dysponuje się danymi można zastosować subiektywne modele tendencji rozwojowej, których parametry nie są szacowane metodami statystycznymi, lecz są określane na podstawie opinii ekspertów [6]. Ze względu na przewidywany kształt krzywej rozwoju rozpatrywanego zjawiska można zastosować trend liniowy, wykładniczy, wykładniczy odwrotnościowy (z asymptotą poziomą) oraz logistyczny. W pracy do budowy prognoz zastosowano trend wykładniczy odwrotnościowy postaci:

$$Y_t = \alpha - \beta g^t, \quad g < 1 \quad (1)$$

W celu uzyskania ocen parametrów $(\hat{\alpha}, \hat{\beta}, \hat{g})$ eksperci muszą określić wartości poziomu analizowanej zmiennej (np. sprzedaży nowego produktu, ilości produkowanej energii z określonej nowej technologii) w pierwszym okresie istnienia na rynku (Y_0), w jednym z późniejszych okresów (Y_n) oraz poziomu nasycenia rynku (Y_∞) (por. tab. 1 w części empirycznej pracy). Oceny parametrów modelu (1) wyznacza się ze wzorów [6]:

* Nr projektu POIG.01.01.01-00-007/08 – Ministerstwo Nauki i Szkolnictwa Wyższego.

$$\hat{\alpha} = \hat{y}_{\infty}, \quad \hat{\beta} = {}^{n-1}\sqrt{\frac{(\hat{y}_{\infty} - \hat{y}_1)^n}{\hat{y}_{\infty} - \hat{y}_n}}, \quad \hat{g} = {}^{n-1}\sqrt{\frac{\hat{y}_{\infty} - \hat{y}_n}{\hat{y}_{\infty} - \hat{y}_1}} \quad (2)$$

gdzie:

$\hat{y}_1, \hat{y}_n, \hat{y}_{\infty}$ – wartości oczekiwane (prognozy punktowe) zmiennych Y_1, Y_n, Y_{∞} .

W przypadku, gdy korzysta się z opinii grupy ekspertów prognozy punktowe mogą być wyznaczone jako średnie arytmetyczne, a gdy korzysta się z sądów pojedynczych ekspertów – wartości oczekiwane rozkładów prawdopodobieństwa subiektywnego*. Prognozę punktową otrzymuje się poprzez ekstrapolację modelu zgodnie z regułą prognozowania według wartości oczekiwanej [6]:

$$y_T^* = f(T, y_1, y_n, y_{\infty}) \quad (3)$$

Dla wyznaczonych prognoz można wyznaczyć standardową niepewność prognoz (odpowiednik błędu prognozy ex ante) $u(y_T)$. Przyjmując, że Y_1, Y_n, Y_{∞} są niezależnymi zmiennymi losowymi, standardowy błąd prognozy można wyznaczyć z formuły [3]:

$$u(y_T) = \left\{ \left[\frac{\partial f}{\partial y_1} u(y_1) \right]^2 + \left[\frac{\partial f}{\partial y_n} u(y_n) \right]^2 + \left[\frac{\partial f}{\partial y_{\infty}} u(y_{\infty}) \right]^2 \right\}^{\frac{1}{2}} \quad (4)$$

gdzie:

$u(y_1), u(y_n), u(y_{\infty})$ – oceny odchyłeń standardowych zmiennych Y_1, Y_n, Y_{∞} ,

$\frac{\partial f}{\partial y_1}, \frac{\partial f}{\partial y_n}, \frac{\partial f}{\partial y_{\infty}}$ – pochodne cząstkowe (współczynniki wrażliwości) funkcji f względem zmiennych odpowiednio Y_1, Y_n, Y_{∞} , liczone w punkcie $t = T, Y_1 = y_1, \dots, Y_n = y_n, Y_{\infty} = y_{\infty}$.

Oceny odchyłeń standardowych $u(y_1), u(y_n), u(y_{\infty})$ można wyznaczyć na podstawie rozkładów prawdopodobieństwa subiektywnego. Występujące we wzorze (4) pochodne cząstkowe dla modelu (1) mają postać [6]:

* Rozkłady prawdopodobieństwa subiektywnego zostały opisane w pracy [3].

$$\begin{aligned}
 \frac{\partial f}{\partial y_1} &= \frac{n-T}{n-1} \left(\frac{\hat{y}_\infty - \hat{y}_n}{\hat{y}_\infty - \hat{y}_1} \right)^{\frac{T-1}{n-1}} \\
 \frac{\partial f}{\partial y_n} &= \frac{T-1}{n-1} \left(\frac{\hat{y}_\infty - \hat{y}_1}{\hat{y}_\infty - \hat{y}_n} \right)^{\frac{n-T}{n-1}} \\
 \frac{\partial f}{\partial y_\infty} &= 1 - \frac{\partial f}{\partial y_1} - \frac{\partial f}{\partial y_n}
 \end{aligned} \tag{5}$$

2. Sposoby konstrukcji prognoz przedziałowych

W pozyskiwaniu danych od ekspertów dotyczących zmian otoczenia i ich wpływu na rozwój danego zjawiska najwygodniejszy jest pomiar zmiennych na skali porządkowej. Można zastosować pięciostopniową skalę semantyczną zaproponowaną przez Osgooda, Suciego i Tannenbauma, gdzie 1 to najniższy poziom czynnika, natomiast 5 to najwyższy poziom czynnika [4]. Dla wyszczególnionych czynników określa się końce skal, np. wolny rozwój – szybki rozwój. Następnie każdemu alternatywnemu poziomowi czynnika kluczowego, dla momentów Y_0 , Y_n i Y_∞ należy przypisać prawdopodobieństwo wystąpienia danego poziomu oraz jego wpływ na rozwój badanego zjawiska (por. tab. 2 i 3 w części empirycznej rozdziału).

Liczba możliwych różnych kombinacji czynników kluczowych wynosi:

$$L = (k + 1)^{n \cdot t} \tag{6}$$

gdzie:

- k – liczba możliwych do przyjęcia alternatywnych wartości w stosunku do wartości bazowej,
- n – liczba wyszczególnionych czynników kluczowych,
- t – liczba okresów, dla których wyznacza się zmiany.

W foresightach często rozpatrywane jest kilkanaście, a nawet kilkadziesiąt czynników kluczowych opisujących otoczenie badanego zjawiska*. W związku z tym, ze względu na ilość rozpatrywanych czynników kluczowych oraz możliwych wartości alternatywnych w stosunku do wartości bazowej czynnika, nie jest możliwe rozpatrzenie wszystkich wariantów prognoz. Wynika to z faktu, że wraz ze wzrostem liczby poziomów alternatywnych oraz liczby rozpatrywanych czynników liczba kombinacji rośnie wykładniczo.

W celu otrzymania przedziałów prognoz, zamiast rozpatrywania wszystkich możliwości kombinatorycznych można zastosować następujące drogi postępowania:

- budowa przedziału prognoz w oparciu o standardową niepewność prognozy,
- agregacja czynników kluczowych w ramach grup tematycznych lub uwzględniając zależność pomiędzy poszczególnymi czynnikami (niekoniecznie z tych samych grup tematycznych),
- symulacyjne wyznaczenie przedziału prognoz na podstawie prawdopodobieństw wystąpienia poziomów czynników kluczowych podanych przez ekspertów.

Pierwsza metoda polega na wyznaczeniu symetrycznego względem prognozy przedziału w oparciu o standardową niepewność prognozy omówionej według formuły [3]:

$$[y_T^* - k_p \cdot u(y_T); y_T^* + k_p \cdot u(y_T)] \quad (7)$$

gdzie:

$u(y_T)$ – standardowa niepewność prognozy y_T ,

k_p – współczynnik związany z wiarygodnością prognozy oraz rozkładem zmiennej prognozowanej wyznaczony z nierówności Czebyszewa.

Jako standardową niepewność prognozy można przyjąć odchylenia standardowe rozkładów trójkątnych. W tym celu eksperci dodatkowo muszą podać dla momentów Y_0 , Y_n i Y_∞ wartość minimalną i maksymalną prognozowanego zjawiska. Omawiane podejście wymaga założenia, że podane przez ekspertów minimalne i maksymalne wartości produkcji/oszczędności energii w pełni uwzględniają zmiany poziomów czynników kluczowych. W przypadku, gdy rozpiętość pomiędzy wartością minimalną i maksymalną będzie zbyt niska/wysoka w stosunku do rzeczywistego wpływu zmian otoczenia, otrzymane przedziały

* W foresighcie „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050” realizowanego w Głównym Instytucie Górnictwa (nr projektu POIG.01.01.01-00-007/08 – Ośrodek Przetwarzania Informacji) uwzględniono czynniki z następujących grup tematycznych: ekonomiczne (29 czynników), społeczne (30 czynników), polityczno-prawne (30 czynników), środowiskowe (17 czynników).

prognoz mogą się okazać zbyt wąskie/szerokie. Uzyskane w ten sposób przedziały są symetryczne względem postawionych prognoz. Jest to istotna wada prezentowanej metody, gdyż zwykle eksperci w swoich sądach częściej wskazują czynniki powodujące wzrost produkcji/oszczędności energii niż jej spadek.

Druga metoda polega na agregacji czynników w ramach grup tematycznych (np. ekonomicznych, społecznych, politycznych, prawnych, środowiskowych). W przypadku, gdy wpływ zmian poziomów czynników kluczowych został podany względnie wobec poziomu bazowego, do agregacji należy użyć ważonej średniej geometrycznej, w której wagami są prawdopodobieństwa (szanse) wystąpienia alternatywnych poziomów czynników kluczowych:

$$\ln \bar{X}_j = \frac{\sum_{i=1}^m w_i \ln x_i}{\sum_{i=1}^m w_i} \quad (8)$$

gdzie:

- $\bar{X}_j$ – średnia ważona czynników z j -tego panelu tematycznego,
- w_i – prawdopodobieństwa wystąpienia alternatywnych poziomów czynników kluczowych,
- x_i – zmiany w oszczędności/produkcji energii spowodowane zmianami poziomów czynników.

W przypadku, gdy wpływ zmian otoczenia na prognozowaną zmienną będzie podany bezwzględnie (w jednostkach naturalnych, np. GWh, sztuki) w stosunku do poziomu bazowego, można użyć ważonej średniej arytmetycznej. Następnie dla wszystkich możliwych kombinacji wyznacza się alternatywne krzywe, a przedział prognozy otrzymujemy na podstawie odpowiednich kwantyli rozkładu prawdopodobieństw danego scenariusza. Wadą tego podejścia jest uśrednienie zmian wszystkich czynników, co wyklucza najbardziej skrajne wartości w poziomie oszczędności/produkcji energii. W efekcie zbudowany przedział prognoz będzie węższy niż w przypadku rozpatrywania wszystkich czynników z osobna.

Kolejna metoda opiera się na symulacyjnym wyznaczeniu przedziałów prognoz. Proces symulacyjny można ująć w czterech etapach:

1. Losowanie poziomów wszystkich czynników kluczowych na podstawie prawdopodobieństw ich wystąpienia określonych przez ekspertów (dane zawarte w tab. 1).

2. Wyznaczenie zagregowanego wpływu zmian poziomów wylosowanych czynników kluczowych w stosunku do wartości bazowej (dane zawarte w tab. 2).
3. Korekta poziomów analizowanej zmiennej dla momentów Y_0 , Y_n i Y_∞ .
4. Wyznaczenie nowych prognoz poprzez ekstrapolację modelu wyznaczonego na podstawie skorygowanych danych.

Postępowanie powtarzane jest n -krotnie, a jego wyniki wyznaczają nowe scenariusze rozwoju analizowanego zjawiska. W celu otrzymania prognoz przedziałowych, dla każdego okresu prognozy wyznacza się miary pozycyjne, np. kwantyl rzędu 0,05 i 0,95. Jako że uwzględnia się w ten sposób także najbardziej skrajne wartości analizowanej zmiennej, błąd wynikający z faktu nieuwzględniania wszystkich możliwych kombinacji będzie mniejszy niż w przypadkach budowania przedziałów prognoz w oparciu o niepewność standardową, czy agregacji czynników w ramach grup tematycznych. Wadą tego podejścia jest stabilność wyników, która zależy od liczby powtórzeń całego procesu oraz empirycznych rozkładów poziomów czynników kluczowych.

3. Wyniki empiryczne

Metody konstrukcji długoterminowych prognoz przedziałowych przedstawiono na podstawie technologii „pompy ciepła”, która została zaliczona do technologii oszczędności zużycia energii finalnej. W tab. 1, 2, 3 przedstawiono część danych pozyskanych od ekspertów.

Tabela 1

Określenie oszczędności produkcji energii finalnej w GWh
dla technologii „pomp ciepła”

Model tendencji rozwojowej	2010 r.			2020 r.			Poziom nasycenia		
	prawd.	max	min	prawd.	max	min	prawd.	max	min
Wykładniczo- odwrotnościowy	200	250	180	750	900	700	900	1000	800

Tabela 2

Przykładowe czynniki otoczenia z panelu ekonomicznego dla technologii „pomp ciepła” wraz z szansami wystąpienia poszczególnych poziomów dla 2050 r.

Nazwa czynnika	Prawdopodobieństwa wystąpienia poziomu czynnika w %				
	1	2	3	4	5
Tempo wzrostu PKB wyższe niż w UE15	6,3	11,3	24,2	35,4	22,8
Wzrost zamożności społeczeństwa	7,3	10,8	24,4	34,1	23,3
Wzrost udziału usług w strukturze PKB	10,3	15,1	23,6	30,8	20,2
Dobra sytuacja finansowa przedsiębiorstw	9,3	13,6	22,5	31,8	22,8
...	...	...	...	...	...

Źródło: Opracowanie własne na podstawie wyników badania „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”.

Tabela 3

Wpływ zmian poziomów przykładowych czynników kluczowych na oszczędności zużycia energii finalnej dla technologii „pomp ciepła” w 2050 r.

Nazwa czynnika	Bazowy poziom czynnika 2050 r.	Zmiana poziomu oszczędności energii w %			
		pierwszy alternatywny poziom czynnika	zmiana	drugi alternatywny poziom czynnika	zmiana
Wzrost zamożności społeczeństwa	4	3	-5	5	10
Wprowadzenie konkurencyjnego rynku energii	4	5	8	3	-10
Wyższe nakłady na edukację w zakresie energooszczędności	4	5	5	3	-5
Mechanizmy i polityka zachęcające do stosowania technologii energoszczędnych	4	5	10	3	-10
...	...	...	...	...	...

Źródło: Opracowanie własne na podstawie wyników badania „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”.

W celu ograniczenia wyników badań do najbardziej prawdopodobnych scenariuszy rozwoju otoczenia, liczbę możliwych poziomów czynników ograniczono do trzech (najbardziej prawdopodobnego oraz dwóch wariantów wyboru). Na ich podstawie otrzymano trzypunktowy rozkład dyskretny*. Pozwoliło to na zredukowanie ilości informacji pozyskiwanych od ekspertów. Prawdopodobieństwa wystąpienia alternatywnych poziomów czynników kluczowych (prawdopodobieństwa z tab. 2) są niezbędne do oceny wag wpływu zmiany poziomu czynnika na produkcję (oszczędność) energii oraz dystrybucyjności empirycznej, wykorzystywanej w symulacji. W tab. 4 przedstawiono zagregowane zmiany w poziomie oszczędności energii z j -tej technologii, wraz z prawdopodobieństwami ich wystąpienia.

Tabela 4

Zagregowane zmiany w poziomie oszczędności energii (%) dla technologii „pompy ciepła” ze względu na alternatywne poziomy czynników kluczowych dla lat 2020 i 2050

Panel tematyczny	2010 r.		2050 r.	
	1. wariant wyboru	2. wariant wyboru	1. wariant wyboru	2. wariant wyboru
Ekonomiczny	5,2 (0,31)	-3,1 (0,22)	5,3 (0,39)	-4,9 (0,16)
Spółeczny	-2,9 (0,35)	0,3 (0,19)	-1,7 (0,41)	-0,4 (0,12)
Polityczno-prawny	-2,4 (0,33)	6,1 (0,22)	-2,3 (0,29)	5,9 (0,16)
Środowiskowy	0,1 (0,22)	-0,1 (0,12)	-0,1 (0,32)	0,1 (0,11)

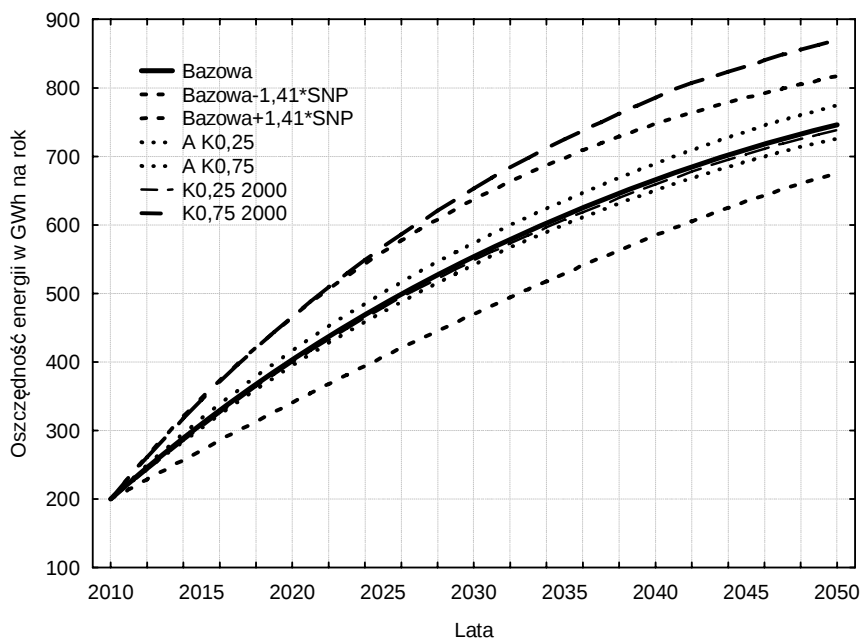
W nawiasach podano prawdopodobieństwa wystąpienia danego alternatywnego poziomu zagregowanego czynnika.

Źródło: Obliczenia własne na podstawie wyników badania „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”.

Na rys. 1 przedstawiono przedziały prognoz zbudowane na podstawie standardowej niepewności prognoz przy współczynniku $k_p = 1,41$ (wartość ta wynika z nierówności Czebyszewa dla $p = 0,50$), przedziały otrzymane poprzez agregację czynników kluczowych oraz na podstawie symulacji (przyjęto kwantyle 025 i 0,75). Na rys. 2 zaprezentowano przedziały prognoz dla wiarygodności $p = 0,90^{**}$.

* Aby trzy najwyższe wartości prawdopodobieństw wystąpienia poziomów czynnika utworzyły rozkład trzypunktowy konieczne było ich przeskalowanie. Polegało ono na proporcjonalnym rozdzieleniu pomiędzy wartości początkowe dopełnienia ich sumy do 100%.

** Z nierówności Czebyszewa przyjęto współczynnik rozszerzenia $k = 3,16$, a w przypadku agregacji i symulacji przyjęto kwantyle 0,05 i 0,95.

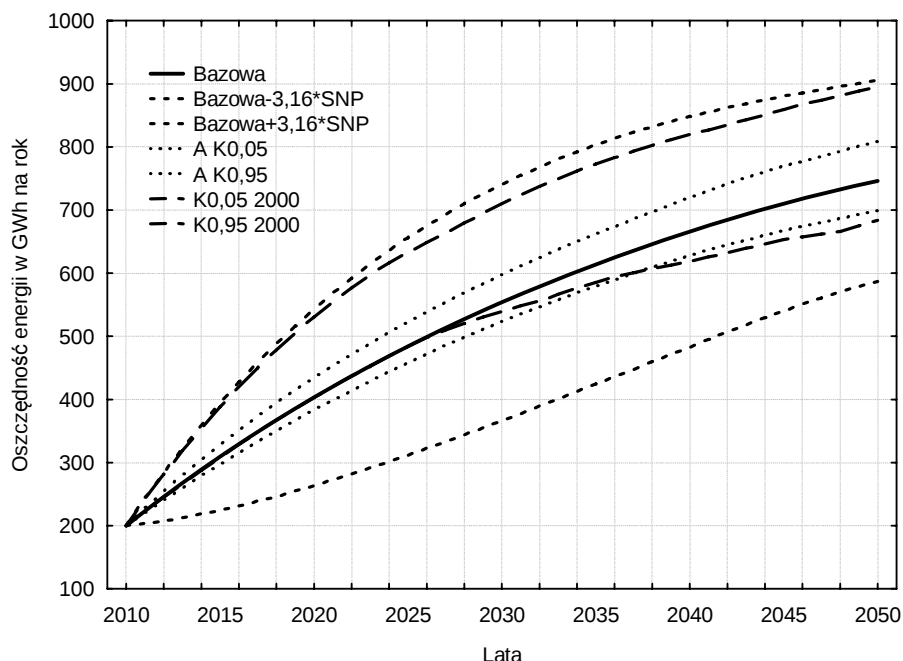


Bazowa $\pm 1,41 \cdot \text{SNP}$ – przedziały prognoz zbudowane w oparciu o standardową niepewność prognoz.

A K0,25; A K0,75 – przedziały prognoz zbudowane za pomocą agregacji czynników kluczowych.
K0,25 2000; K0,75 2000 – przedziały prognoz zbudowane za pomocą symulacji wyznaczone na podstawie 2000 replikacji.

Rys. 1. Prognozy bazowe wraz z przedziałami prognoz wyznaczonymi na podstawie kwantyla 0,25 i 0,75

Źródło: Opracowanie własne na podstawie wyników badania „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”.



Bazowa $\pm$ 3,16*SNP – przedziały prognoz zbudowane w oparciu o standardową niepewność prognoz.

A K0,05; A K0,95 – przedziały prognoz zbudowane za pomocą agregacji czynników kluczowych.

K0,05 2000; K0,95 2000 – przedziały prognoz zbudowane za pomocą symulacji wyznaczone na podstawie 2000 replikacji.

Rys. 2. Prognozy bazowe wraz z przedziałami prognoz wyznaczonymi na podstawie kwantyla 0,05 i 0,95

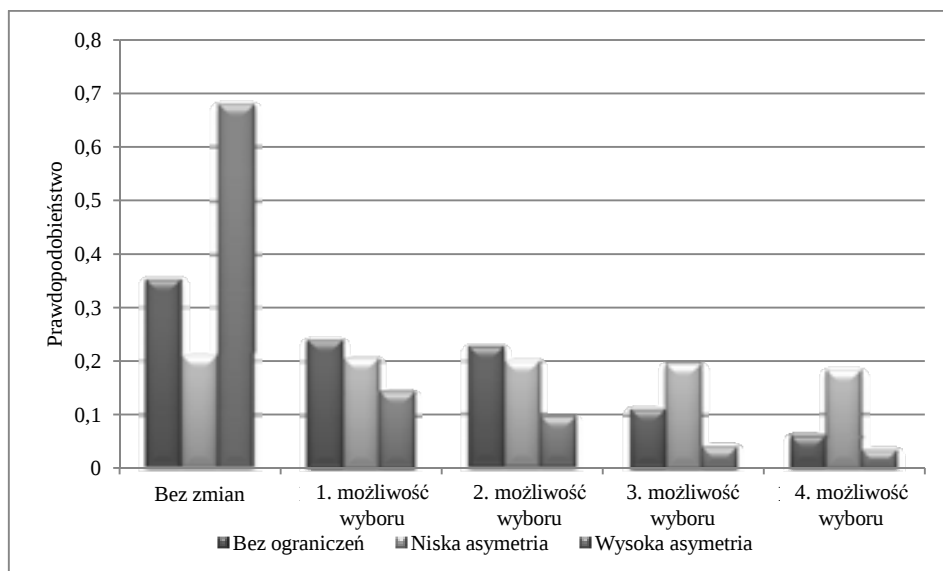
Źródło: Opracowanie własne na podstawie wyników badania „Zeroemisyjna gospodarka energią w warunkach zrównoważonego rozwoju Polski do 2050”.

W przypadku prognoz wyznaczonych z 50% i 90% prawdopodobieństwem najszerzy przedział prognoz uzyskano dla przedziałów zbudowanych w oparciu o standardową niepewność prognoz, natomiast najwęższy dla agregacji czynników*. Przedziały prognoz uzyskane na podstawie zmian czynników kluczowych były asymetryczne względem prognoz bazowych, gdyż eksperci częściej wskazywali na determinanty rozwoju technologii „pomp ciepła” niż na bariery.

W następnym etapie analizy porównano zbieżność wyników dla wyznaczonych granic przedziałów prognoz. W tym celu wygenerowano rozkłady o różnej asymetrii dla pięciu i trzech poziomów czynników kluczowych. Przyjęto dwa skrajne rozkłady: o niewielkiej asymetrii (rozkład zbliżony do rozkładu równo-

* Prawie trzykrotnie węższy od najszerzego.

miernego – drugi słupek na rys. 3) oraz rozkład skrajnie asymetryczny (prawdopodobieństwo dla bazowego poziomu czynnika jest bardzo wysokie, a dla kolejnych wariantów wyboru jest niskie – trzeci słupek na rys. 3). Dodatkowo wygenerowano rozkład, w którym nie założono żadnych ograniczeń (pierwszy słupek na rys. 3).



Rys. 3. Przykładowe rozkłady prawdopodobieństw wystąpienia poziomów czynników kluczowych

Tabela 5

Odchylenia standardowe granic przedziałów prognoz wyznaczonych na podstawie 100 replikacji.

Rozkład o silnej asymetrii – kwantyl 0,95				
1	2	3	4	5
Liczba losowań	2020	2030	2040	2050
100	5,187	8,389	10,618	12,267
200	4,081	6,608	8,557	10,095
500	4,390	6,699	8,541	10,230
1000	2,413	3,003	3,738	3,642
2000	2,309	3,185	3,468	3,251
5000	1,161	2,317	3,438	2,384

cd. tabeli 5

1	2	3	4	5
Liczba losowań	2020	2030	2040	2050
Rozkład o słabej asymetrii – kwantyl 0,75				
Liczba losowań	2020	2030	2040	2050
100	0,372	0,542	1,992	4,410
200	0,414	0,570	0,628	0,624
500	0,976	2,463	4,189	5,735
1000	0,000	0,000	0,000	0,000
2000	0,000	0,000	0,000	0,000
5000	0,000	0,000	0,000	0,000
Rozkład bez ograniczeń – kwantyl 0,05				
Liczba losowań	2020	2030	2040	2050
100	3,424	5,175	5,954	6,353
200	2,113	3,150	3,725	4,267
500	2,129	2,746	2,868	3,088
1000	1,170	1,383	0,842	0,567
2000	0,947	0,920	0,270	0,019
5000	0,909	0,898	0,248	0,019

Wyznaczono przedziały prognoz na podstawie: 0,25 i 0,75 kwantyla oraz 0,05 i 0,95 kwantyla dla 100, 200, 500, 1000, 2000, 5000 losowań. Każdą liczbę losowań powtórzono 100 razy i na ich podstawie wyznaczono: medianę, średnią arytmetyczną oraz odchylenie standardowe. W tab. 5 przedstawiono wyniki dla poszczególnych pięciopunktowych (wartość bazowa i 4 poziomy alternatywne) rozkładów prawdopodobieństw czynników kluczowych. Uwzględniono tam kwantyl, w którym wystąpiła najwolniejsza zbieżność mierzona odchyleniem standardowym.

Najszybciej zbieżne są wyniki dla rozkładu o słabej asymetrii, w którym 1000 losowań wystarczy, aby wszystkie replikacje wyznaczyły te same granice prognoz. W przypadku rozkładu o najsilniejszej asymetrii przy 1000 losowań można zauważyć gwałtowny spadek wartości odchylen standardowych, jednak dopiero 5000 losowań gwarantowało stabilność otrzymanej granicy przedziału prognoz (współczynnik zmienności nie przekraczał 5%). W przypadku rozkładu bez żadnych ograniczeń współczynnik zmienności dla kwantyla 0,05 nie przekraczał 2,5%. Stąd można wysnuć wniosek, że 5000 iteracji zapewnia zbieżność na wysokim poziomie. Gdy założono, że czynniki kluczowe przyjmują trzy poziomy, zbieżność symulacji była szybsza. W przypadku każdego rozkładu 2000 losowań gwarantowało uzyskanie takich samych wyników dla każdej replikacji.

Wnioski

Budowa długoterminowych prognoz rozwoju nowych zjawisk powinna uwzględniać zmienność otoczenia we wskazanym w badaniu horyzoncie. Gdy zmiany otoczenia są systematyczne i powolne można w tym celu wykorzystać prognozy wariantowe. Wadą tego podejścia jest gwałtownie rosnąca liczba prognoz wraz ze wzrostem liczby zmiennych opisujących otoczenie lub rozpatrywanych alternatywnych ich poziomów. Znacznie mniej skomplikowanym numerycznie sposobem jest budowa prognoz przedziałowych. Wśród omawianych metod na uwagę zasługuje metoda wyznaczania przedziałów na podstawie symulacji. W odróżnieniu od innych metod nie odrzuca ona nawet najbardziej skrajnych wartości analizowanej zmiennej, przez co uwzględnia nietypowe scenariusze, a jednocześnie charakteryzuje się dość szybką zbieżnością. Do jej wad można zaliczyć dużą liczbę danych, jakie należy uzyskać od ekspertów, co może niekorzystnie wpłynąć na precyzję ich sądów, a tym samym na wiarygodność zbudowanych prognoz. Pewnego rodzaju rozwiązaniem tego problemu może być zredukowanie liczby cech opisujących otoczenie lub liczby alternatywnych poziomów. Przedziały prognoz otrzymane na podstawie agregacji czynników uśredniają zmiany otoczenia, co w efekcie daje znacznie węższy przedział prognoz niż w przypadku analizy symulacyjnej, w związku z tym jego rozpiętość nie jest doszacowana. Przedział prognozy oparty na standardowej niepewności prognoz jest symetryczny względem prognoz bazowych. Nie uwzględnia tym samym jednej z istotnych informacji pozyskanych od ekspertów, jaką jest asymetria wpływu zmian otoczenia na prognozowane zjawisko.

Literatura

1. Armstrong J.S., Collopy F., *Integration of Statistical Methods and Judgment for Time Series Forecasting*, w: G. Wright, P. Goodwin, *Forecasting with Judgment*, John Wiley & Sons, New York 1998.
2. Dittmann P., *Prognozowanie w przedsiębiorstwie. Metody i ich zastosowanie*, Oficyna Ekonomiczna, Kraków 2004.
3. Gogolewska (Poradowska) K., *Ocena dopuszczalności prognoz gospodarczych*, Wrocław 2006 (praca doktorska).
4. *Metody statystycznej analizy wielowymiarowej w badaniach marketingowych*, red. E. Gatnar, M. Walesiak, Akademia Ekonomiczna, Wrocław 2004.
5. Pawłowski Z., *Prognozy ekonometryczne*, PWN, Warszawa 1973.
6. Poradowska K., *Wybrane aspekty prognozowania wielkości sprzedaży nowych produktów*, w: *Modelowanie i prognozowanie gospodarki narodowej*, Prace i Materiały Wydziału Zarządzania Uniwersytetu Gdańskiego nr 5/2007, Sopot 2007.

THE SELECTED ASPECTS OF CONSTRUCTING LONG-TERM INTERVAL FORECASTS FOR THE DEVELOPMENT OF NEW PHENOMENA

Summary

With regard to a long horizon of the forecasts built, among others, for the needs of foresight deep changes should be expected in the area of the considered phenomenon. In this connection, the variability of the surrounding, i.e., economic, political, legal, social and technological situation as well as the natural environment should be taken into account in the process of creating a long-term forecast. This problem can be partially solved by means of constructing variant forecasts that take into consideration the adopted scenarios of the surrounding development. However, in such a case, we obtain merely point forecasts that, from the point of view of the construction of development scenarios, may be insufficient. The analysis should be enriched by means of forecast intervals, which would take into account the changes in the development of the considered phenomenon with regard to the change in the level of key factors. The purpose of the article is the presentation of the way to build long-term point forecasts that would take into consideration the changes in the surrounding of the analysed phenomenon. The author compares the intervals of forecasts obtained by means of standard forecast uncertainty, key factors aggregation and simulation analysis.

Alicja Wolny-Dominiak

Uniwersytet Ekonomiczny w Katowicach

MODELOWANIE LICZBY SZKÓD W UBEZPIECZENIACH KOMUNIKACYJNYCH W PRZYPADKU WYSTĘPOWANIA DUŻEJ LICZBY ZER, Z WYKORZYSTANIEM PROCEDURY KROSWALIDACJI

Wprowadzenie

Jednym z problemów występujących w analizie danych ubezpieczeniowych jest modelowanie liczby szkód występujących w danym portfelu polis z zastosowaniem regresji przy założeniu rozkładu Poissona. Portfele ubezpieczeniowe charakteryzują się jednak tym, że dla wielu polis w okresie ubezpieczenia nie wystąpiła żadna szkoda. Oznacza to, iż dane zawierają dużą liczbę zer, co powoduje, że klasyczna regresja Poissona nie daje zadowalających wyników. W pierwszej części pracy przedstawiono uogólnioną regresję Poissona dla zmiennej licznikowej oraz zmodyfikowaną wersję regresji Poissona uwzględniającą sytuację występowania dużej liczby zer w danych (*zero-inflated Poisson regression*). W drugiej części przeprowadzono przykład empiryczny możliwości zastosowania wersji zmodyfikowanej do modelowania liczby szkód w ubezpieczeniach majątkowych. Analizowano różne modele w celu określenia, które zmienne taryfikacyjne wpływają na występowanie zer w portfelu polis stosując procedurę 10-krotnej krosvalidacji. W efekcie uzyskano ranking pozwalający na klasyfikację polis ze względu na liczbę generowanych szkód. Dane do przykładu obliczeniowego zaczerpnięto z literatury przedmiotu. Do obliczeń wykorzystano program komputerowy R, pakiet `{pscl}` oraz zaimplementowany algorytm krosvalidacji (załącznik A).

1. Modele regresji z licznikową zmienną objaśnianą typu ZI

W tej części pracy przedstawione są modele regresji, w których zmienną zależną jest zmienna licznikowa przyjmująca wartości całkowite nieujemne oraz występuje duża liczba wartości zerowych (*zero-inflated*). W ubezpieczeniach majątkowych takie modele mają zastosowanie w szczególności w modelowaniu oraz prognozowaniu liczby szkód. Stosowany jest najczęściej model regresji Poissona, w którym przyjmuje się założenie, że zmienna objaśniana Y ma rozkład Poissona $Y \sim \text{Pois}(\lambda)$ warunkowany wartościami zmiennych objaśniających [2]:

$$P(Y_i = y_i) = \frac{e^{-\lambda_i} \lambda_i^{y_i}}{y_i!}, \quad i = 1, \dots, n$$

W powyższym wzorze Y_i oznacza liczbę szkód dla i -tej osoby ubezpieczonej. Parametr λ_i uzależniony jest od pewnych zmiennych zależnych X_j , $j = 1, \dots, k$ charakteryzujących ubezpieczonego oraz pojazd, którego dotyczy ubezpieczenie, np. płci, wieku, pojemności silnika samochodu. Najczęściej przyjmowana jest logarytmiczna funkcja połączenia:

$$\ln \lambda_i = \sum_{j=1}^k \beta_{ji} X_{ji}$$

Korzystając z własności rozkładu Poissona, że parametr λ_i jest równy wartości oczekiwanej, mamy:

$$\lambda_i = \mu_i = e^{\sum_{j=1}^k \beta_{ji} X_{ji}}$$

Widać zatem, że dla każdej kombinacji zmiennych objaśniających uzyskiwana jest zawsze dodatnia oczekiwana liczba szkód. W modelu przyjmuje się założenia, że zmienna Y ma rozkład Poissona, średnia wartość zmiennej jest równa wariancji oraz $y_1, \dots, y_n$ są niezależne o stałej wariancji.

Parametr λ_i może być wykorzystywany do rangowania polis ze względu na liczbę szkód. Niezbędna jest jednak korekta tego parametru wskaźnikiem ekspozycji na ryzyko dla i -tej polisy d_i , który pokazuje najczęściej w przypadku ubezpieczeń majątkowych, jaką część badanego okresu obejmowała polisa:

$$\lambda_i = d_i e^{\beta_0 + \sum_{j=1}^k \beta_{ij}}$$

Powyższy model nie uwzględnia przypadku, w którym zmienna licznikowa przyjmuje dużą liczbę wartości zerowych. Taka sytuacja występuje często w przypadku modelowania liczby szkód. Analizując portfele ryzyk można założyć, że dla wielu polis nie wystąpiła żadna szkoda, natomiast w przypadku wystąpienia szkód są to jedna, dwie, trzy i rzadko więcej. Dlatego w przypadku analizy liczby szkód w zakładzie ubezpieczeń zasadniejsze wydaje się stosowanie zmodyfikowanej regresji Poissona, gdzie uwzględnia się dużą liczbę wartości zerowych w danych zwanej modelem ZIP (*Zero-Inflated Poisson*). W modelu ZIP niezależne zmienne Y_i przyjmują wartości zerowe: $Y_i \sim 0$ z prawdopodobieństwem ϖ_i lub wartości z rozkładu Poissona: $Y_i \sim \text{Pois}(\lambda_i)$ z prawdopodobieństwem $1 - \varpi_i$, co można zapisać następująco [5]:

$$P(Y_i = y_i) = \begin{cases} \varpi_i + (1 - \varpi_i)e^{-\lambda_i}, & y_i = 0 \\ (1 - \varpi_i) \frac{e^{(-\lambda_i)\lambda_i^{y_i}}}{y_i!}, & y_i > 0 \end{cases}, \quad i = 1, \dots, n$$

Zatem w modelu ZIP występują dwa parametry: λ_i oraz ϖ_i . Oba parametry, podobnie jak w przypadku regresji Poissona, połączone są ze zmiennymi objaśniającymi następującymi funkcjami połączeń:

$$\ln\left(\frac{\varpi_i}{1 - \varpi_i}\right) = \sum_{j=1}^t \gamma_{ji} Z_{ji}$$

$$\ln \lambda_i = \sum_{j=1}^k \beta_{ji} X_{ji},$$

gdzie $Z_1, \dots, Z_t$ są zmiennymi zależnymi dla równania pierwszego, natomiast $X_1, \dots, X_k$ zmiennymi dla równania drugiego.

Oczekiwana liczba szkód oraz wariancja liczby szkód i -tej polisy w modelu ZIP wynosi [5]:

$$E(Y_i) = \lambda_i(1 - \varpi_i)$$

$$E(Y_i) = (1 - \varpi_i)(\lambda_i - \varpi_i \lambda_i^2)$$

Podobnie jak w przypadku regresji Poissona, w modelu ZIP zakłada się, iż średnia liczba szkód jest równa wariancji. W przypadku, gdy wariancja jest wyższa od średniej występuje problem nadmiernej dyspersji, który często charakteryzuje zmienne licznikowe. Powoduje on, że statystyki χ^2 testujące istotność parametrów strukturalnych modelu są przeszacowane, natomiast nie

zmienia zgodności estymatorów parametrów. W celu uniknięcia nadmiernej dyspersji można zastosować skorygowane błędy standardowe lub przejść do modelu, w którym wprowadzany jest rozkład negatywny-dwumodalny [4].

Do wyboru modelu szacowania liczby szkód ubezpieczeniowych oraz układu zmiennych wpływających na generowanie przez polisy wartości zero-owych możliwe jest wykorzystanie koncepcji statystycznych metod automatycznego uczenia się. Ogólnie w tej koncepcji zakłada się, że dany jest zbiór uczący $D = \{(x^i, y^i), i = 1, \dots, N\}$, gdzie $x^i, y^i \in R$. Ponadto zbiór uczący tworzą obserwacje wylosowane z jednakowym prawdopodobieństwem w sposób niezależny, z populacji o wielowymiarowym rozkładzie określonym przez nieznaną funkcję gęstości:

$$p(x, y) = p(x)p(y | x)$$

Zadanie polega na przeszukaniu pewnego zbioru funkcji $H = \{f(x, \varpi) : \varpi \in \Omega\}$, gdzie ϖ jest wektorem parametrów modelu, i wskazaniu elementu najlepszego. Posługując się modelem $f(x, \varpi) \in H$, który jest uproszczonym obrazem analizowanego zjawiska, w trakcie przeszukiwania pojawiają się błędy wynikające z wykorzystywania wartości teoretycznych w miejscu wartości rzeczywistych zmiennej objaśnianej. Błędy mierzone są tzw. funkcjami straty $L(y, f(y, \varpi))$, które najczęściej mierzą błąd predykcji dla pojedynczej obserwacji. W koncepcji metod automatycznego uczenia się rozważany jest całkowity błąd modelu będący sumą wartości funkcji straty dla wszystkich możliwych obserwacji. Jedną z metod estymacji wartości błędu całkowitego jest metoda sprawdzania krzyżowego (*CV-cross-validation*) [3]. W niniejszej pracy zastosowano następujący algorytm:

- a) losowe wyznaczenie ze zbioru danych 10 podzbiorów o zbliżonej liczebności, $k = 10$, (n – liczebność całego zbioru, m_l – liczebność l -tego podzbioru, $l = 1, \dots, 10$),
- b) 10-krotne szacowanie modelu na podzbiorze danych o liczebności $n - m_l$ z usunięciem zbioru walidującego,

$$c) \text{ 10-krotne wyznaczenie błędu } MSE_l = \frac{\sum (y - \hat{\mu}_l)^2}{m_l},$$

$$d) \text{ szacowanie błędu krosvalidacji: } cv = \sum_{l=1}^{10} \frac{m_l}{n} MSE_l.$$

Porównując modele wybrano model o najmniejszej wartości cv . Implementację procedury w programie komputerowym R przedstawiono w załączniku A.

2. Przykład empiryczny

Proces modelowania i prognozowania liczby szkód w zakładzie ubezpieczeń przeprowadzono z wykorzystaniem bazy danych szkód komunikacyjnych (*third party motor insurance claims*) zaczerpniętej z pozycji [1]. Baza danych zawiera następujące zmienne uwzględnione w modelu:

1. Zmienna objaśniana licznikowa:
numclaims – liczba szkód.
2. Zmienne objaśniające:
veh_body – kształt samochodu,
veh_age – wiek samochodu: A (najmłodszy), B, C, D,
gender – płeć kierowcy: M (kobieta), F,
agecat – wiek kierowcy: A (najmłodszy), B, C, D, E, F.

Obliczenia wykonano w programie komputerowym R. Rozkład liczby szkód w analizowanym portfelu przedstawia się następująco:

Tabela 1

Rozkład liczby szkód

Liczba szkód	Liczba szkód	Częstość	Średnia ekspozycja na ryzyko
0	63232	93,19%	0,45
1	4333	6,39%	0,6
2	271	0,40%	0,71
3	18	0,03%	0,7
4	2	0,00%	0,88

Jak widać, liczba szkód charakteryzuje się bardzo dużą liczbą zer, gdzie 93% polis nie wygenerowało żadnej szkody w portfelu. Wartość wariancji przewyższa wartość średniej i indeks nadmiernej dyspersji jest na poziomie:

$$O = \frac{\text{wariancja} - \text{średnia}}{\text{średnia}} = 0,0063$$

co oznacza słaby efekt nadmiernej dyspersji w portfelu. Do modelowania liczby szkód zastosowano w pierwszej kolejności regresję Poissona. W modelu P1 badano wpływ poszczególnych zmiennych na liczbę szkód:

$$\ln \lambda = \beta_0 + \beta_1 \text{veh_body} + \beta_2 \text{veh_age} + \beta_3 \text{gender} + \beta_4 \text{agecat}$$

Model szacowano wykorzystując funkcję `glm(){stats}`, przyjmując rozkład Poissona dla liczby szkód. W pierwszej kolejności zbadano istotność wpływu poszczególnych zmiennych na liczbę szkód.

Tabela 2

Parametry strukturalne regresji Poissona dla modelu P1

Model P1	$\hat{\beta}_i$	Średni błąd szacunku	p-wartość
Stała	-1,18	0,32	0,00
<i>Veh_body</i>	-0,95	0,39	0,05
<i>Veh_age</i>	-0,04	0,01	0,00
<i>Gender</i>	-0,01	0,03	0,79
<i>Agecat</i>	-0,08	0,01	0,00

W modelu P1 na poziomie istotności 5% zmienna charakteryzująca płeć jest statystycznie nieistotna, dlatego w dalszej analizie zmienna ta została usunięta z modelu, pozostałe zmienne nie są skorelowane. Nowy model P2 przyjął postać:

$$\ln \lambda = \beta_0^1 + \beta_1^1 \text{veh_body} + \beta_2^1 \text{veh_age} + \beta_3^1 \text{agecat}$$

Uzyskane parametry strukturalne zamieszczono w tab. 3.

Tabela 3

Parametry strukturalne regresji Poissona dla modelu P2

Realizacje zmiennych w modelu P2	β	e^β	Średni błąd szacunku
1	2	3	4
Stała	-1,35	0,26	0,32
<i>veh_ageA</i>	0,00	1,00	–
<i>veh_ageB</i>	0,13	1,14	0,04
<i>veh_ageC</i>	0,001	1,001	0,04
<i>veh_ageD</i>	-0,08	0,93	0,04
<i>AgecatA</i>	0,00	1,00	–
<i>AgecatB</i>	-0,17	0,85	0,05
<i>AgecatC</i>	-0,20	0,82	0,05
<i>AgecatD</i>	-0,23	0,80	0,05
<i>AgecatE</i>	-0,42	0,65	0,06

cd. tabeli 3

1	2	3	4
<i>AgecatF</i>	-0,43	0,65	0,07
<i>veh_body_BUS</i>	1,00	1,00	–
<i>veh_body_CONVT</i>	-1,75	0,17	0,66
<i>veh_body_COUPE</i>	-0,75	0,47	0,34
<i>veh_body_HBACK</i>	-1,10	0,33	0,32
<i>veh_body_HDTOP</i>	-0,87	0,42	0,33
<i>veh_body_MCARA</i>	-0,46	0,63	0,41
<i>veh_body_MIBUS</i>	-1,15	0,32	0,35
<i>veh_body_PANVN</i>	-0,84	0,43	0,34
<i>veh_body_RDSTR</i>	-0,68	0,51	0,66
<i>veh_body_SEDAN</i>	-1,04	0,35	0,32
<i>veh_body_STNWG</i>	-1,00	0,37	0,32
<i>veh_body_TRUCK</i>	-1,04	0,35	0,33
<i>veh_body_UTE</i>	-1,25	0,29	0,32

W estymacji parametrów strukturalnych modelu przyjęto zmienne bazowe jako: *veh_ageA*, *AgecatA*, *veh_body_BUS*. Interpretując uzyskane wyniki, na podstawie wartości parametrów strukturalnych zawartych w tab. 3 można stwierdzić kierunek wpływu zmiany wieku samochodu, wieku kierowcy oraz kształtu samochodu na liczbę szkód na podstawie znaku. Widać więc, że przy ustalonym układzie zmiennych bazowych, stopy tariff będą obniżały składkę. W celu określenia jednostkowego wpływu zmiennych objaśniających na liczbę szkód wyznaczono eksponenty parametrów strukturalnych modelu (w modelu przyjęto logarytmiczną funkcję połączenia).

Wszystkie parametry modelu są statystycznie istotne. Również test ilorazu wiarygodności pokazał, że model jest w całości statystycznie istotny. W wyniku działania funkcji `lrtest{lmtest}` uzyskano bardzo niski, prawie zerowy poziom p-wartości.

Tabela 4

Test ilorazu wiarygodności dla modelu P2

	#Df	LogLik	Df	Chisq	Pr(> Chisq)
Model P1.1	21	-18029,2	NA	NA	NA
Model – tylko stała	1	-18101,5	-20	144,5839	0,000001

Do rankingu polis w modelu P2 wykorzystano parametr λ . Minimalna wartość tego parametru wyniosła $\hat{\lambda} = 0,0271$ i jest to kategoria polis generująca najmniejszą liczbę szkód: (*veh_bodyCONVT*, *veh_ageD*, *agecatF*). Klasa polis generująca największą liczbę szkód uzyskała $\hat{\lambda} = 1,2712$ dla kategorii (*veh_bodyRDSTR*, *veh_ageB*, *agecatA*).

Ponieważ w analizowanej bazie danych występuje duża liczba polis, dla których nie wystąpiła żadna szkoda, dalej dokonano szacowania różnorodnych modeli ZIP analizując wpływ różnych zmiennych taryfikacyjnych na wystąpienie dużej liczby zer. Do wyboru i ostatecznej postaci modelu zastosowano procedurę 10-krotnej krosvalidacji.

Model ZIP1

$$\ln \lambda^{ZIP1} = \beta_0^{ZIP1} + \beta_1^{ZIP1} \text{veh_body} + \beta_2^{ZIP1} \text{veh_age} + \beta_3^{ZIP1} \text{agecat}$$

$$\ln\left(\frac{\varpi^{ZIP1}}{1 - \varpi^{ZIP1}}\right) = \gamma_0^{ZIP1}$$

Model ZIP2

$$\ln \lambda^{ZIP2} = \beta_0^{ZIP2} + \beta_1^{ZIP2} \text{veh_body} + \beta_2^{ZIP2} \text{veh_age} + \beta_3^{ZIP2} \text{agecat}$$

$$\ln\left(\frac{\varpi^{ZIP2}}{1 - \varpi^{ZIP2}}\right) = \gamma_0^{ZIP2} + \gamma_1^{ZIP2} \text{veh_body}$$

Model ZIP3

$$\ln \lambda^{ZIP3} = \beta_0^{ZIP3} + \beta_1^{ZIP3} \text{veh_body} + \beta_2^{ZIP3} \text{veh_age} + \beta_3^{ZIP3} \text{agecat}$$

$$\ln\left(\frac{\varpi^{ZIP3}}{1 - \varpi^{ZIP3}}\right) = \gamma_0^{ZIP3} + \gamma_1^{ZIP3} \text{veh_body} + \gamma_2^{ZIP3} \text{veh_age}$$

Model ZIP4

$$\ln \lambda^{ZIP4} = \beta_0^{ZIP4} + \beta_1^{ZIP4} \text{veh_body} + \beta_2^{ZIP4} \text{veh_age} + \beta_3^{ZIP4} \text{agecat}$$

$$\ln\left(\frac{\varpi^{ZIP5}}{1 - \varpi^{ZIP5}}\right) = \gamma_0^{ZIP5} + \gamma_1^{ZIP5} \text{veh_body} + \gamma_2^{ZIP5} \text{agecat}$$

Model ZIP5

$$\ln \lambda^{ZIP5} = \beta_0^{ZIP5} + \beta_1^{ZIP5} veh_body + \beta_2^{ZIP5} veh_age + \beta_3^{ZIP5} agecat$$

$$\ln\left(\frac{\varpi^{ZIP5}}{1-\varpi^{ZIP5}}\right) = \gamma_0^{ZIP5} + \gamma_1^{ZIP5} veh_age + \gamma_2^{ZIP5} agecat$$

Model ZIP6

$$\ln \lambda^{ZIP6} = \beta_0^{ZIP6} + \beta_1^{ZIP6} veh_body + \beta_2^{ZIP6} veh_age + \beta_3^{ZIP6} agecat$$

$$\ln\left(\frac{\varpi^{ZIP6}}{1-\varpi^{ZIP6}}\right) = \gamma_0^{ZIP6} + \gamma_2^{ZIP6} veh_age$$

Model ZIP7

$$\ln \lambda^{ZIP7} = \beta_0^{ZIP7} + \beta_1^{ZIP7} veh_body + \beta_2^{ZIP7} veh_age + \beta_3^{ZIP7} agecat$$

$$\ln\left(\frac{\varpi^{ZIP7}}{1-\varpi^{ZIP7}}\right) = \gamma_0^{ZIP7} + \gamma_3^{ZIP7} agecat$$

Modele szacowano wykorzystując funkcję `zeroinfl(){pscl}` oraz procedurę krosvalidacji zaimplementowaną w programie komputerowym R. Kod programu zawarto w załączniku A. Uzyskano następujące całkowite błędy cv dla analizowanych modeli:

- $cv^1 = 0,077296$
- $cv^2 = 0,071232$
- $cv^3 = 0,077186$
- $cv^4 = 0,073221$
- $cv^5 = 0,076565$
- $cv^6 = 0,077112$
- $cv^7 = 0,076651$

Widać zatem, że modelem, który daje najmniejszy błąd cv jest model, gdzie na wystąpienie zera wpływa zmienna veh_body ZIP2.

Tabela 5

Parametry strukturalne modelu ZIP2

Realizacje zmiennych w modelu P2	β	e^{β}	Średni błąd szacunku
Stała	-1,27	0,28	1,15
<i>veh_ageA</i>	0,00	1,00	–
<i>veh_ageB</i>	0,13	1,14	0,04
<i>veh_ageC</i>	0,00	1,00	0,04
<i>veh_ageD</i>	-0,08	0,93	0,05
<i>AgecatA</i>	0,00	1,00	–
<i>AgecatB</i>	-0,17	0,85	0,06
<i>AgecatC</i>	-0,20	0,82	0,05
<i>AgecatD</i>	-0,23	0,80	0,05
<i>AgecatE</i>	-0,42	0,65	0,06
<i>AgecatF</i>	-0,43	0,65	0,07
<i>veh_body_BUS</i>	1,00	1,00	–
<i>veh_body_CONVT</i>	-1,83	0,16	1,34
<i>veh_body_COUPE</i>	-0,13	0,88	1,21
<i>veh_body_HBACK</i>	-0,84	0,43	1,16
<i>veh_body_HDTOP</i>	-0,96	0,38	1,15
<i>veh_body_MCARA</i>	-0,54	0,58	1,31
<i>veh_body_MIBUS</i>	-0,87	0,42	1,35
<i>veh_body_PANVN</i>	-0,24	0,79	1,22
<i>veh_body_RDSTR</i>	1,38	3,96	1,48
<i>veh_body_SEDAN</i>	-0,34	0,71	1,15
<i>veh_body_STNWG</i>	-0,61	0,55	1,16
<i>veh_body_TRUCK</i>	-0,37	0,69	1,19
<i>veh_body_UTE</i>	-0,66	0,52	1,18

Porównując uzyskane wyniki w modelu P2 oraz modelu ZIP2 uwzględniającym występowanie dużej liczby zer w modelu widać, że parametry dla zmiennej *veh_body* znacznie się różnią w obu modelach, natomiast pozostałe parametry są niezmiennie. Wynika to z faktu, iż w procedurze krosvalidacji uzyskano wynik pokazujący, że na generowane przez polisy wartości zerowe wpływa jedynie kształt samochodu.

Do rankingu polis w modelu ZIP2 wykorzystano parametr λ . Minimalna wartość tego parametru wyniosła $\tilde{\lambda} = 0,0271$ i jest to kategoria polis generująca najmniejszą liczbę szkód: (*veh_bodyCONVT*, *veh_ageD*, *agecatF*). Kategoria polis generująca największą liczbę szkód uzyskała $\tilde{\lambda} = 0,2952$ i jest to kategoria (*veh_bodyBUS*, *veh_ageB*, *agecatA*). Szczegółowy ranking zawarto w załączniku B.

Podsumowanie

W pracy przedstawiono zmodyfikowaną regresję Poissona w przypadku, gdy w danych występuje duża liczba zer dla zmiennej licznikowej i jej porównanie z klasyczną regresją Poissona. Zaproponowano zastosowanie k -krotnej krosvalidacji do wyboru czynników wpływających na generowanie przez polisy zerowych liczb szkód. Dodatkowo wyznaczając odpowiednie parametry rozkładu stworzono ranking polis według kategorii zmiennych taryfikacyjnych. Zastosowanie rankingu w praktyce pozwala na sklasyfikowanie nowo zawieranej polisy do odpowiedniej grupy taryfikacyjnej. Zasadniczą wadą klasycznej regresji Poissona, jak również modeli ZIP jest fakt, iż w danej klasie polis wszystkie polisy charakteryzują się taką samą oczekiwaną liczbą szkód, co jest założeniem mało realnym. Rozwiązaniem tego problemu jest przejście do mieszanego modelu Poissona wprowadzając czynnik losowy różnicujący polisy.

Załącznik A

Implementacja procedury krosvalidacji w programie komputerowym R

```

library(psc1)
car=read.csv2(file="c:/car.csv")
K=10
options(OutDec="," )
mse.cv=function(dataset, model, K=10){
  cvseg=c()
  set.seed(113)
  cvseg=cvsegments(nrow(dataset), K)
  ModelMSE=c()
  for (i in 1:K) {
    validset=NULL
    validset=eval(parse(text=paste("cvseg$V", i, sep="")))
    datasetTrainCV= NULL;
    datasetValidCV= NULL
    datasetTrainCV= dataset[-validset,]
    datasetValidCV= dataset[validset,]
    Formula=model$formula
    Model.na.cv=NULL
    Model.na.cv=zeroinfl(formula = Formula, data=datasetTrainCV)
    datasetValidCV.bez.y=NULL
    datasetValidCV.bez.y= subset(datasetValidCV, select=c(veh_body,
veh_age, agecat))
    pred.valid=NULL
    pred.valid=predict(Model.na.cv, datasetValidCV.bez.y)
    MSE.valid=NULL
    MSE.valid=sum((datasetValidCV$numclaims-
pred.valid)^2)/length(pred.valid)
    ModelMSE=c(ModelMSE, MSE.valid)
  }
  MSE.CV=NULL
  MSE.CV=mean(ModelMSE)
  mse.cv.glm=function(dataset, model, K=10){
    cvseg=c()
    set.seed(113)
    cvseg=cvsegments(nrow(dataset), K)
    ModelMSE=c()
    for (i in 1:K) {
      validset=NULL
      validset=eval(parse(text=paste("cvseg$V", i, sep="")))
      datasetTrainCV= NULL; datasetValidCV= NULL

```

```

datasetTrainCV= dataset[-validset,]
datasetValidCV= dataset[validset,]
Formula=model$formula
Model.na.cv=NULL
Model.na.cv=glm(formula = Formula, family=poisson,
data=datasetTrainCV)
datasetValidCV.bez.y=NULL
datasetValidCV.bez.y= subset(datasetValidCV, select=c(veh_body,
veh_age, agecat))
pred.valid=NULL
pred.valid=predict(Model.na.cv, datasetValidCV.bez.y)
MSE.valid=NULL
MSE.valid=sum((datasetValidCV$numclaims-
pred.valid)^2)/length(pred.valid)
ModelMSE=c(ModelMSE, MSE.valid)
}
MSE.CV=NULL
MSE.CV=mean(ModelMSE)

g=glm(formula=numclaims ~ veh_body+veh_age+agecat, family=poisson, data=car)
z1=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | 1, data = car)
z1.mse=NULL
z1.mse=mse.cv(dataset=car, model=z1, K=10)
z2=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | veh_body, data = car)
z2.mse=NULL
z2.mse=mse.cv(dataset=car, model=z3, K=10)
z3=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | veh_body+veh_age,
data = car)
z3.mse=NULL
z3.mse=mse.cv(dataset=car, model=z4, K=10)
z4=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | veh_body+agecat,
data = car)
z4.mse=NULL
z4.mse=mse.cv(dataset=car, model=z5, K=10)
z5=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | veh_age+agecat, da-
ta = car)
z5.mse=NULL
z5.mse=mse.cv(dataset=car, model=z6, K=10)
z6=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | veh_age, data = car)
z6.mse=NULL
z6.mse=mse.cv(dataset=car, model=z7, K=10)
z7=zeroinfl(formula = numclaims ~ veh_body+veh_age+agecat | agecat, data = car)
z7.mse=NULL
z7.mse=mse.cv(dataset=car, model=z8, K=10)

```

Załącznik B

Ranking polis ze względu na kategorię zmiennych taryfikacyjnych w modelu ZIP2

Kształt samochodu	Wiek samochodu	Wiek kierowcy	$\tilde{\lambda}$
veh_bodySEDAN	veh_ageD	agecatF	0,0550
veh_bodySEDAN	veh_ageD	agecatE	0,0556
veh_bodySEDAN	veh_ageA	agecatF	0,0596
veh_bodySEDAN	veh_ageC	agecatF	0,0596
veh_bodySEDAN	veh_ageA	agecatE	0,0602
veh_bodySEDAN	veh_ageC	agecatE	0,0602
veh_bodySEDAN	veh_ageD	agecatD	0,0672
veh_bodySEDAN	veh_ageB	agecatF	0,0679
veh_bodySEDAN	veh_ageB	agecatE	0,0686
veh_bodySEDAN	veh_ageD	agecatC	0,0693
veh_bodySEDAN	veh_ageD	agecatB	0,0714
veh_bodySEDAN	veh_ageA	agecatD	0,0728
veh_bodySEDAN	veh_ageC	agecatD	0,0728
veh_bodySEDAN	veh_ageA	agecatC	0,0750
veh_bodySEDAN	veh_ageC	agecatC	0,0750
veh_bodySEDAN	veh_ageA	agecatB	0,0773
veh_bodySEDAN	veh_ageC	agecatB	0,0773
veh_bodySEDAN	veh_ageB	agecatD	0,0829
veh_bodySEDAN	veh_ageD	agecatA	0,0846
veh_bodySEDAN	veh_ageB	agecatC	0,0854
veh_bodySEDAN	veh_ageB	agecatB	0,0880
veh_bodySEDAN	veh_ageA	agecatA	0,0916
veh_bodySEDAN	veh_ageC	agecatA	0,0916
veh_bodySEDAN	veh_ageB	agecatA	0,1044

Literatura

1. de Jong P., Heller G.Z., *Generalized Linear Models for Insurance Data*, Cambridge University Press 2008.
2. Denuit M., Marechal X., Pitrebois S., Walhin J., *Actuarial Modelling of Claims Counts*, John Wiley & Sons Ltd, 2007.
3. Gatnar E., *Podejście wielomodelowe w zagadnieniach dyskryminacji i regresji*, Wydawnictwo Naukowe PWN, Warszawa 2008.
4. Kopczewska K., Kopczewski T., Wójcik P., *Metody ilościowe w R. Aplikacje ekonomiczne i finansowe*, CeDeWu, Warszawa 2009.
5. Lambert D., *Zero-Inflated Poisson Regression, with an Application to Defects in Manufacturing*, „Technometrics” 1992, Vol. 34, No. 1.

MODELING THE NUMBER OF CLAIMS IN MOTOR INSURANCE IN CASE OF A LARGE NUMBER OF ZEROS USING THE PATCH VALIDATION PROCEDURES

Summary

The problem in the analysis of insurance data is modeling the number of claims occurring in a given portfolio policy using regression assuming a Poisson distribution which is not always justified, since sometimes the data contains a large number of zeros. This paper presents a generalized Poisson regression for the counter variable and a modified version of Poisson regression taking into account the situation of the presence of a large number of zeros in the data (called zero-inflated Poisson regression). Various types were analyzed in order to determine which variables influence the occurrence of zeros in the portfolio policy using the procedure 10 times the patch validation. The result is ranking for classification policies because of the number of generated damage.

Tadeusz Czernik

Uniwersytet Ekonomiczny w Katowicach

WPŁYW NIEPEWNOŚCI OSZACOWANIA ZMIENNOŚCI NA CENĘ INSTRUMENTÓW POCHODNYCH

Wprowadzenie

Jednym z filarów współczesnych finansów jest teoria wyceny instrumentów pochodnych. Spośród wielu modeli wyceny model Blacka-Scholesa [1; 8] w literaturze przedmiotu wyróżnia się prostotą i olbrzymią liczbą odwołań. W większości opracowań poruszana jest kwestia wyceny waniliowych i egzotycznych opcji. Spora liczba prac wykazuje nieadekwatność modelu Blacka-Scholesa (efekt „uśmiechu” zmienności) [5; 7]. W kręgu zainteresowań badaczy i praktyków znajduje się również analiza wrażliwości ceny instrumentu na zmianę parametrów modelu. Losowa natura aktualnych cen derywatów nie jest następstwem stochastycznej ewolucji cen akcji, lecz jest indukowana statystycznym charakterem procedury estymacji parametrów modelu. Poniższe rozważania dotyczą budowy przedziałów ufności aktualnych cen wybranych instrumentów pochodnych.

1. Model Blacka-Scholesa

Rynek instrumentów pochodnych zawdzięcza swój rozkwit m.in. modelowi wyceny instrumentów pochodnych Blacka-Scholesa (w 1997 r. Myron Scholes wraz z Robertem Mertonem za badania dotyczące wyceny derywatów otrzymali nagrodę Nobla z ekonomii; Fischer Black zmarł w 1995 r.). Pomimo wad model ten jest jednym z najczęściej stosowanych modeli wyceny.

Założenia modelu:

- 1) inwestor w każdej chwili może zakupić dowolną (nawet niecałkowitą) ilość akcji (instrument bazowy) i instrumentu wolnego od ryzyka,
- 2) dozwolona jest krótka sprzedaż,

- 3) brak kosztów transakcyjnych (w tym podatków),
- 4) inwestor może pożyczać i lokować kapitał po tej samej stałej stopie (intensywności) wolnej od ryzyka,
- 5) brak dywidendy,
- 6) ewolucja cen akcji jest geometrycznym ruchem Browna (stochastyczne równanie różniczkowe rozumiane jest w sensie Ito) [3; 6]:

$$dS_t = \mu S_t dt + \sigma S_t dW_t \quad (1)$$

gdzie:

μ – dryf,

σ – zmienność,

S_t – cena akcji,

W_t – proces Wienera,

- 7) na rynku (instrument wolny od ryzyka, akcja, instrument pochodny) nie występuje arbitraż, tzn. nie można osiągnąć zysku bez narażania się na ryzyko.

Pierwsze założenie nie jest realistyczne, gdyż nie można zakupić/sprzedać niecałkowitej liczby akcji. Założenia tego nie możemy opuścić, gdyż w wyprowadzeniu równania Blacka-Scholesa „zaszyta” jest koncepcja replikacji, tzn. konstrukcja portfela imitującego wyceniany instrument. W praktyce inwestorzy posiadają w swoich portfelach wiele identycznych derywatów, z których każdy opiewa np. na zakup lub sprzedaż 100 akcji. Skutkiem tego jest duża (najprawdopodobniej niecałkowita) liczba akcji w portfelu replikującym. Ponieważ niemożliwy jest zakup/sprzedaż niecałkowitej liczby akcji, pożądana liczba akcji w portfelu replikującym jest zaokrąglana do liczby całkowitej. Podczas tej fazy wprowadza się błąd rzędu ułamka promila, co sprawia że powyższe założenie możemy zaakceptować.

W pierwszym założeniu ukryte jest również założenie doskonałej płynności rynku, co oznacza równowagę popytu i podaży.

Drugie założenie w zasadzie zawiera się w pierwszym, jednak ze względu na jego istotność autorzy postanowili je wyeksponować.

Trzecie założenie odgrywa istotną rolę, gdyż wraz z upływem czasu zmienia się skład portfela replikującego (dynamiczny hedging). Uwzględnienie kosztów transakcyjnych sprawia, że doskonała replikacja nie jest możliwa. Zagadnienie wyceny w otoczeniu ekonomicznym z kosztami transakcyjnymi sprowadza się do rozwiązania zagadnienia optymalnego sterowania lub po zastosowaniu pewnych aproksymacji do nieliniowego cząstkowego równania różniczkowego. Ponieważ w przypadku „dużych” portfeli inwestorzy mogą wynegocjować niskie koszty transakcyjne, założenie nie wprowadza istotnego odstępstwa od realiów rynkowych.

Czwarte założenie składa się *de facto* z dwóch założeń: istnieje instrument wolny od ryzyka oraz stopa oprocentowania tego instrumentu jest jednakowa dla lokaty i pożyczki (cena „zakupu” i cena „sprzedaży” jest identyczna). Przyjmuje się, że za instrument pozbawiony ryzyka można przyjąć instrumenty rządowe (emitowane przez skarbu państwa). Wydarzenia ostatnich kilku lat sprawiły, że do powyższej reguły powinniśmy podchodzić z rezerwą (Islandia, Grecja, Włochy). Należy jednak pamiętać, że na ryzyko tych papierów wpływają czynniki makroekonomiczne, a sygnały o pogarszającej się sytuacji emitenta są identyfikowalne z wyprzedzeniem. Oznacza to, że jeżeli parametry makroekonomiczne emitenta nie wskazują na jego złą kondycję, możemy założyć, że wyemitowany przez niego instrument będzie praktycznie pozbawiony ryzyka w okresie uzależnionym od charakteru przeprowadzonej analizy makroekonomicznej.

Równość stóp wynika przede wszystkim z pierwszego założenia (płynność/równowaga) oraz faktu, iż jest to instrument pozbawiony ryzyka.

Piąte założenie stanowi istotne odstępstwo od realiów rynkowych, nie jest to jednak argument za odrzuceniem modelu. Omawiana tu wersja modelu wyceńny jest jedną z najprostszych. Uwzględnienie dywidendy nie stanowi dużego problemu (pozostaje problem modelowania przyszłej dywidendy).

Szóste założenie wynika po pierwsze z faktu, iż w przypadku niewielkiej niepewności przyszłej ceny ($\sigma \ll 1$) ewolucja ceny powinna być zgodna z zasadą kapitalizacji ciągłej, po drugie, żądanie stacjonarności infinitezymalnych stóp zwrotu $\frac{dS_t}{S_t} = \frac{S_{t+dt} - S_t}{S_t}$ znacznie upraszcza rozważania/obliczenia. Po

trzecie, proces Wienera był i jest nadal jednym z najlepiej poznanych i spularyzowanych procesów stochastycznych. Oczywiście kierowano się również zgodnością z rzeczywistą dynamiką cen (rozkład stóp zwrotu jest rozkładem logarymiczno-normalnym). Należy podkreślić, że skonstruowano również modele lepiej odzwierciedlające dynamikę cen instrumentu bazowego oraz ceny derywatów. W niniejszym opracowaniu założono, że parametry modelu stopa wolna od ryzyka, dryf oraz zmienność są stałe, jednak nic nie stoi na przeszkodzie, aby modelować je deterministycznymi funkcjami czasu lub procesami losowymi. Warto podkreślić, że z szóstego założenia wynika również, że podmiot, który wyemitował akcję nie może zbankrutować, gdyż cena akcji jest zawsze większa od zera – rozwiązanie równania (1) jest większe od zera:

$$S_t = S_0 e^{\left(\mu - \frac{1}{2}\sigma^2\right)t + \sigma W_t} > 0 \quad (2)$$

Założenie braku arbitrażu jest jednym z filarów metod wyceny. Sformułowanie: „nie można osiągnąć ponadprzeciętnych zysków bez narażania się na ryzyko” oznacza, że np. w sytuacji, gdy początkowa wartość nakładów jest większa od zera, wartość stopy zwrotu nie może być tylko większa lub równa od stopy instrumentu wolnego od ryzyka, przy czym prawdopodobieństwo, że stopa zwrotu z inwestycji w portfel będzie większa od stopy wolnej od ryzyka jest dodatnie. Podobnie jak większość omówionych wcześniej założeń nie jest ono założeniem realistycznym. W realiach rynkowych występuje arbitraż, a graczy wykorzystujących możliwość osiągnięcia zysku bez narażania się na ryzyko nazywa się arbitrażystami. Niemniej jednak w dobie powszechnej informatyzacji sytuacji, w których można osiągnąć zysk ze strategii arbitrażowej występują stosunkowo krótko. Mechanizm podaży-popytu sprawia, iż arbitraż zanika bardzo szybko. Stąd stwierdzamy, że założenie to jest akceptowalne.

Ponieważ jedynymi zmiennymi są czas t i cena instrumentu bazowego S , możemy założyć, że cena derywatu będzie funkcją czasu i ceny akcji $V(S, t)$. Z lematu Ito [6] wynika, że ewolucja ceny instrumentu pochodnego będzie dana następującym stochastycznym równaniem różniczkowym (pominięto argumenty (S, t)):

$$dV = \left(\frac{\partial V}{\partial t} + \mu S \frac{\partial V}{\partial S} + \frac{1}{2} \sigma^2 S^2 \frac{\partial^2 V}{\partial S^2} \right) dt + \sigma S \frac{\partial V}{\partial S} dW \quad (3)$$

gdzie dW jest identyczny z przyrostem procesu Wienera występującym we wzorze (1). Oznacza to możliwość skonstruowania portfela pozbawionego niepewności, czyli również ryzyka (niepewność jest warunkiem koniecznym wystąpienia ryzyka).

Załóżmy, że wystawiliśmy opcję i zamierzamy zneutralizować ryzyko krótkiej pozycji (hedging dynamiczny) dokonując zakupu pewnej liczby akcji Δ . Ewolucja wartości tak skonstruowanego portfela $P = \Delta S - V$ dana jest wzorem (wymagamy również, aby portfel był portfelem samofinansującym, w przeciwnym przypadku występowałby instrument nieuwzględniony na naszym rynku):

$$\begin{aligned} dP &= \Delta dS - dV = \\ &= \left(-\frac{\partial V}{\partial t} + \mu S \left(\Delta - \frac{\partial V}{\partial S} \right) - \frac{1}{2} \sigma^2 S^2 \frac{\partial^2 V}{\partial S^2} \right) dt + \sigma S \left(\Delta - \frac{\partial V}{\partial S} \right) dW \end{aligned} \quad (4)$$

Jak łatwo zauważyć, dobierając liczbę akcji Δ tak, aby $\Delta = \frac{\partial V}{\partial S}$ sprawimy, że stochastyczne równanie różniczkowe (4) nie będzie zawierało czynnika odpowiadającego za losowość (ryzyko):

$$dP = \frac{\partial V}{\partial S} dS - dV = \left(-\frac{\partial V}{\partial t} - \frac{1}{2} \sigma^2 S^2 \frac{\partial^2 V}{\partial S^2} \right) dt \quad (5)$$

Ponieważ wartość tak skonstruowanego portfela ewoluuje w sposób deterministyczny (wolny od ryzyka), jego dynamika musi być identyczna z dynamiką wartości instrumentu wolnego od ryzyka. W przeciwnym przypadku można by skonstruować strategię arbitrażową (jedno z założeń wykluczało wystąpienie arbitrażu). Oznacza to, że:

$$dP = rPdt = r \left(\frac{\partial V}{\partial S} S - V \right) dt \quad (6)$$

czyli:

$$r \left(\frac{\partial V}{\partial S} S - V \right) dt = \left(-\frac{\partial V}{\partial t} - \frac{1}{2} \sigma^2 S^2 \frac{\partial^2 V}{\partial S^2} \right) dt \quad (7)$$

Ostatecznie otrzymujemy (dzieląc przez $dt > 0$):

$$\frac{\partial V}{\partial t} + \frac{1}{2} \sigma^2 S^2 \frac{\partial^2 V}{\partial S^2} + rS \frac{\partial V}{\partial S} - rV = 0 \quad (8)$$

Uzupełniając powyższe równanie (Blacka-Scholesa) o odpowiedni warunek (końcowy, brzegowy) i rozwiązując powyższe równanie różniczkowe cząstkowe otrzymamy formułę wyceny instrumenty pochodnego.

Na podkreślenie zasługuje fakt, że aktualna wartość opcji nie zależy od współczynnika dryfu μ , lecz jedynie od stopy wolnej od ryzyka, zmienności, czasu, aktualnej ceny akcji i parametrów pochodzących z warunków końcowych i brzegowych. Oznacza to, że wartość opcji jest wielkością deterministyczną, w pełni determinowaną powyższymi parametrami.

2. Przedziałowe oszacowanie ceny opcji

W praktyce nie znamy rzeczywistych wartości zmienności σ i stopy wolnej od ryzyka r . Niniejsze opracowanie dotyczy jedynie przypadku nieznanego wartości zmienności. W celu jego oszacowania stosuje się narzędzia analizy statystycznej. Ponieważ immanentną cechą procedur statystycznych jest niepewność otrzymanych wielkości, niepewność oszacowania zmienności propaguje się na oszacowanie ceny opcji $V(S, t, \sigma, r, \dots)$, gdzie kropki oznaczają zależność od parametrów występujących w warunkach końcowym oraz brzegowym.

W przypadku europejskich opcji waniliowych *call* i *put* wyznaczenie przedziałów ufności ich cen nie przedstawia trudności, gdyż ich ceny są rosnącymi funkcjami zmienności (dla $t < T$) [4]:

$$Vega = \frac{\partial c}{\partial \sigma} = \frac{\partial p}{\partial \sigma} = SN'(d_1)\sqrt{T-t} > 0 \quad (9)$$

gdzie:

$$d_1 = \frac{\ln\left(\frac{S}{K}\right) + \left(r + \frac{1}{2}\sigma^2\right)(T-t)}{\sigma\sqrt{T-t}},$$

c – wartość opcji *call*,

p – wartość opcji *put*,

K – cena wykonania,

T – data wykonania opcji,

t – aktualny moment czasu,

$N(\cdot)$ – dystrybuenta standardowego rozkładu normalnego ($N'(\cdot)$ – gęstość standardowego rozkładu normalnego).

Oszacowania przedziałowe cen mają w tym przypadku postać (poziom ufności wynosi 0,95; symetryczny podział istotności):

$$(V(S, t, \sigma_{0,025}), V(S, t, \sigma_{0,975})) \quad (10)$$

gdzie:

V – wartość opcji *call/put*,

$$\sigma_{\frac{\alpha}{2}}^2 = \frac{(n-1)\hat{D}^2 \ln\left(\frac{S_{t+\Delta t}}{S_t}\right)}{\Delta t \chi_{1-\frac{\alpha}{2}, n-1}^2},$$

$$\sigma_{1-\frac{\alpha}{2}}^2 = \frac{(n-1)\hat{D}^2 \ln\left(\frac{S_{t+\Delta t}}{S_t}\right)}{\Delta t \chi_{\frac{\alpha}{2}, n-1}^2},$$

$\hat{D}^2 \ln \left(\frac{S_{t+\Delta t}}{S_t} \right)$ – wartość nieobciążonego estymatora wariancji logarytmicznych stóp zwrotu,

n – liczebność próby,

Δt – odstęp czasu między notowaniami (w opracowaniu przyjęto $\Delta t = \frac{1}{250}$).

Warto podkreślić, że wielkości σ_{α}^2 i $\sigma_{1-\alpha}^2$ nie są kwantylami kwadratu

zmienności (zmienność σ jest wielkością deterministyczną), lecz są końcami przedziału, który na zadanym poziomie ufności pokrywa rzeczywistą wartość kwadratu zmienności.

W ogólności wartość instrumentu pochodnego nie jest monotoniczną funkcją zmienności. Można wtedy zastosować wzory transformacyjne dla prawdopodobieństw lub gęstości, wymaga to jednak znajomości punktów, w których zmienia się rodzaj monotoniczności ceny instrumentu pochodnego. Autorzy wykorzystali rozwiązanie symulacyjne. Wygenerowane kwadraty zmienności podstawiano do wzoru na wartość opcji, a następnie wyznaczono empiryczne kwantyle cen.

Przykładowe wykresy zaprezentowano dla europejskiej opcji *call Cash-or-Nothing* (członek większej rodziny opcji typu *binary/digital*). Profil wypłaty tej opcji ma postać [4]:

$$Payoff = \begin{cases} X & \text{gdy } S_T > K \\ 0 & \text{gdy } S_T \leq K \end{cases} \quad (11)$$

gdzie X jest kwotą wypłacaną, gdy opcja jest *in-the-money* ($S_T > K$), S_T jest ceną instrumentu bazowego (akcji) w dniu wykonania opcji.

Cena wyznaczona z modelu Blacka-Scholesa dana jest wzorem [4]:

$$c = Xe^{-r(T-t)} N \left(\frac{\ln \left(\frac{S}{K} \right) + \left(r - \frac{1}{2} \sigma^2 \right) (T-t)}{\sigma \sqrt{T-t}} \right) \quad (12)$$

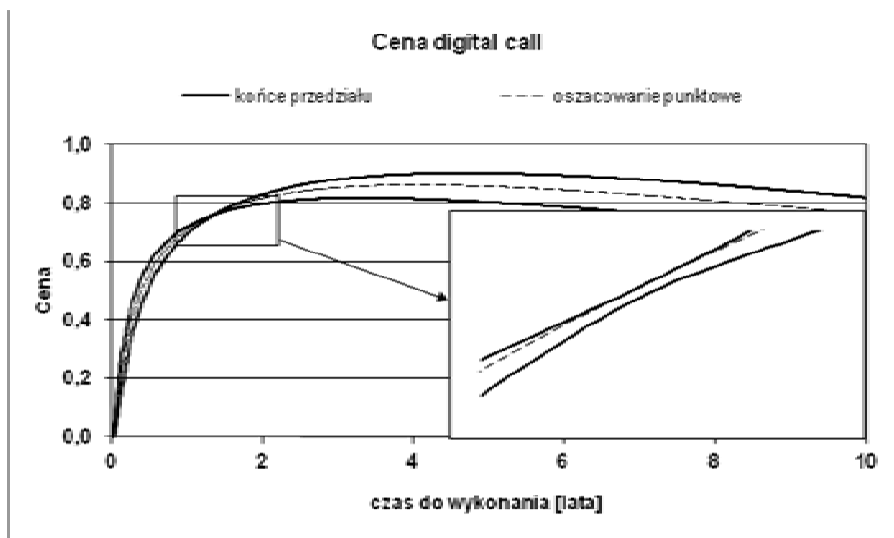
Obliczając pochodną ze względu na zmienność otrzymujemy:

$$\begin{aligned} \text{Vega} &= \frac{\partial c}{\partial \sigma} = \\ &= -Xe^{-r(T-t)}N'\left(\frac{\ln\left(\frac{S}{K}\right) + \left(r - \frac{1}{2}\sigma^2\right)(T-t)}{\sigma\sqrt{T-t}}\right)\left(\frac{\ln\left(\frac{S}{K}\right) + r(T-t)}{\sigma^2} + \frac{1}{2}\sqrt{T-t}\right) \end{aligned} \quad (13)$$

W przypadku, gdy wyrażenie $\ln\left(\frac{S}{K}\right) + r(T-t)$ jest ujemne, cena opcji ma mieszaną monotoniczność (ze względu na zmienność).

Rysunek 1 przedstawia przedziałowe oszacowanie ceny opcji w zależności od czasu do wykonania $T-t$ ($n=100$, $K=1,1$, $X=2$, $S=1$, $r=0,05$,

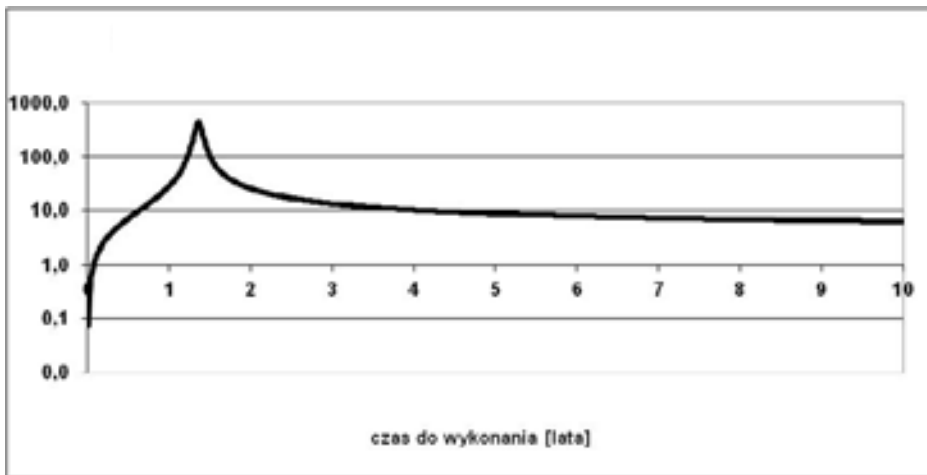
$$\hat{\sigma}^2 = \frac{\hat{D}^2 \ln\left(\frac{S_{t+\Delta t}}{S_t}\right)}{\Delta t} = 0,04).$$



Rys. 1. Przedziałowe oszacowanie ceny europejskiej opcji digital w zależności od czasu wykonania. Wartości pozostałych parametrów w tekście

Zauważmy, że rozpiętość przedziału ufności ma minimum lokalne w okolicach czasu równego około 1,4. W okolicach tego czasu do wykonania znajduje się miejsce zerowe pochodnej (minimum wrażliwości ceny ze względu na zmienność). Gdyby wyznaczyć aproksymację odchylenia standardowego rozkładu cen okazałoby się, że z dokładnością do wyrazów pierwszego rzędu jego wartość wynosi zero. Oznacza to, że w okolicach minimum wrażliwości należy stosować aproksymacje co najmniej drugiego rzędu (proste wprowadzenie do poruszanej wyżej techniki aproksymacji można znaleźć np. w [2]).

Rysunek 2 przedstawia stosunek oszacowania punktowego ceny opcji do rozpiętości przedziału ufności (wartości parametrów jak wyżej).

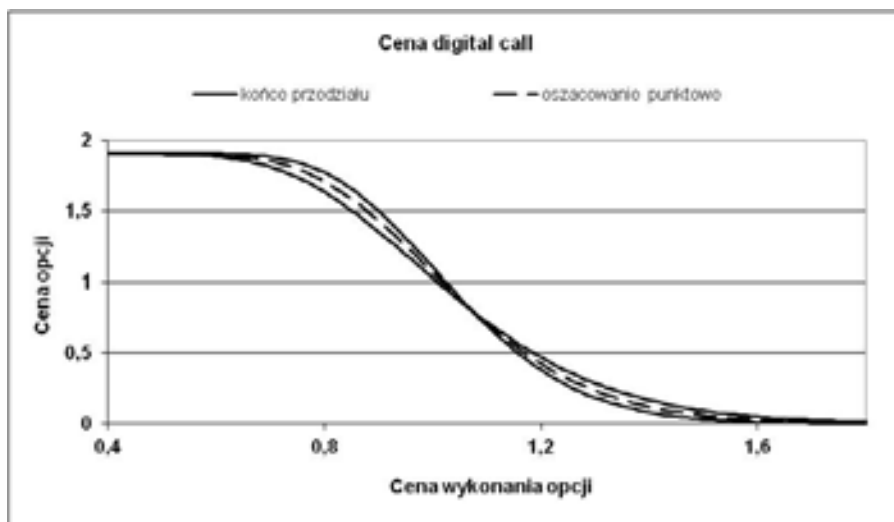


Rys. 2. Stosunek oszacowania punktowego do rozpiętości przedziału ufności w zależności od czasu wykonania. Wartości pozostałych parametrów w tekście

Jak wynika z rys. 2, największą względną dokładność otrzymujemy dla czasu do wykonania wynoszącego około 1,4. Błąd oszacowania mierzony rozpiętością przedziału stanowi ułamek procenta oszacowania punktowego, jednak dla większości czasów błąd ten jest rzędu 10% oszacowania punktowego.

Rysunek 3 przedstawia oszacowanie przedziałowe ceny opcji dla różnych wartości ceny wykonania ($n=100$, $T-t=1$, $X=2$, $S=1$, $r=0,05$,

$$\hat{\sigma}^2 = \frac{\hat{D}^2 \ln\left(\frac{S_{t+\Delta t}}{S_t}\right)}{\Delta t} = 0,04).$$



Rys. 3. Przedziałowe oszacowanie ceny europejskiej opcji digital w zależności od ceny wykonania. Wartości pozostałych parametrów w tekście

Podobnie jak poprzednio, minimum lokalne rozpiętości przedziału osiągnęte jest punkcie, w którym *Vega* opcji jest równa zero ($K \approx 1,07$).

Rysunek 4 przedstawia stosunek punkowego oszacowania do rozpiętości przedziału ufności w zależności od ceny wykonania (parametry, jak w przypadku rys. 3).

Podobnie jak wyżej zauważamy, że najmniejszy względny błąd oszacowania (poniżej jednego procenta) ceny osiągany jest (lokalnie) w okolicach ceny wykonania równej 1,07, jednak w przeważającej części obszaru (wykresu), w którym opcja jest *out-of-the-money* wartość względnego błędu jest rzędu kilkudziesięciu procent.



Rys. 4. Stosunek oszacowania punkowego do rozpiętości przedziału ufności w zależności od ceny wykonania. Wartości pozostałych parametrów w tekście

Jak można było zauważyć w powyższych rozważaniach, rozpiętość oszacowania przedziałowego ceny opcji jest nietrywialną funkcją parametrów rynkowych.

Z uwagi na ograniczoną ilość miejsca nie przedstawiono oszacowań przedziałowych w zależności od pozostałych parametrów.

Podsumowanie

W większości opracowań poruszających wycenę instrumentów pochodnych podawane są tylko wzory analityczne na wartość derywatu. Spory odsetek prac porusza również analizę wrażliwości ceny na zmiany wartości parametrów. Omówiona praca jest jedną z niewielu, w których autorzy przedstawiają propagację błędu oszacowania parametrów stochastycznej ewolucji ceny instrumentu bazowego na cenę derywatu. Przeprowadzona analiza pokazuje, że zależność błędu oszacowania mierzonego rozpiętością przedziału ufności od parametrów modelu jest nietrywialna. Można się spodziewać, że w przypadku opcji o bardziej skomplikowanej funkcji wypłaty (np. w przypadku opcji egzotycznych, koszykowych) oraz w bardziej realistycznych modelach stochastycznej dynamiki złożoność zależności będzie wyższa od przedstawionej w pracy.

Literatura

1. Bjork T., *Arbitrage theory in continuous time*, Oxford University Press 2009.
2. Casella G., Berger R.L., *Statistical inference*, Cengage Learning 2009.
3. Czernik T., *Skazani na formalizm Ito?*, w: *Metody matematyczne, ekonometryczne i informatyczne w finansach i ubezpieczeniach*, red. P. Chrzan, Akademia Ekonomiczna, Katowice 2006.
4. Haug E.G., *The complete option pricing formulas*, McGraw-Hill 2007.
5. Latane H., Rendleman R., *Standard deviations of stock price ratios implied in option prices*, „J. Finance” 1976, No. 31.
6. Oksendal B., *Stochastic differential equations*, Springer 2007.
7. Rubinstein M., *Implied Binomial Trees*, „Journal of Finance” 1994, Vol. 49, No. 3.
8. Shreve S.E., *Stochastic calculus for finance. Continuous-time model*, Springer 2008.
9. Sobczyk M., *Statystyka*, PWN, Warszawa 2002.

ASSESS THE IMPACT OF UNCERTAINTY ON THE PRICE VOLATILITY DERIVATIVES

Summary

Valuation of derivatives is one of the most discussed topics of scientific treatises. In this paper we assess the likely impact of uncertainty on the price volatility of derivative. Results are presented on the example of the European digital option. It has been shown non-trivial dependence of the span of the confidence interval of the model parameters.

Monika Dyduch

Uniwersytet Ekonomiczny w Katowicach

BANKOWE PAPIERY WARTOŚCIOWE STRUKTURYZOWANE

Wprowadzenie

Bankowe Papiery Wartościowe Strukturyzowane (BPW) emitowane* są na podstawie ustawy z 29 sierpnia 1997 r. Prawo bankowe (tekst jednolity: Dz.U. 2002, nr 72 poz. 665, z późn. zm.). Emitent zobowiązuje się względem posiadaczy BPW do spełnienia świadczenia pieniężnego polegającego na zapłacie kwoty wykupu oraz premii na zasadach określonych w dokumentach regulujących warunki emisji danego BPW**.

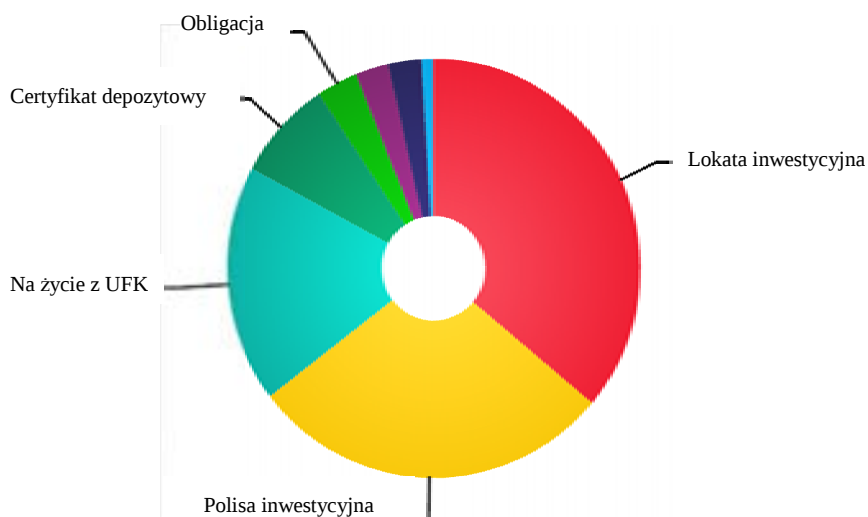
Treść BPW obejmuje:

- wartość nominalną,
- zobowiązanie banku do:
 - naliczenia określonego oprocentowania według ustalonej stopy procentowej,
 - dokonania wypłaty oznaczonej kwoty osobie uprawnionej, w określonych terminach; osoba uprawniona nie może żądać od banku wykupu papieru przed upływem terminu o ile treść papieru nie stanowi inaczej,
- oznaczenie posiadacza papieru wartościowego, jeżeli jest to papier imienny lub adnotację, że jest to papier wartościowy na okaziciela,
- zasady przenoszenia praw wynikających z papieru wartościowego,
- numer papieru wartościowego i datę emisji,
- podpisy osób upoważnionych do składania oświadczeń w zakresie praw i obowiązków majątkowych banku.

* W treści BPW, jak również w podanej przez emitenta do publicznej wiadomości informacji o warunkach emisji nie mogą być zamieszczane porównania z warunkami emisji papierów wartościowych innych emitentów.

** Bank nie może udzielać kredytu lub pożyczki pieniężnej na kupno BPW emitowanych przez siebie.

Opracowanie przedstawia przykładową wycenę BPW strukturyzowanego, czyli produktu strukturyzowanego w formie prawnej, która stanowiła 3,21% proponowanych w okresie 27.11.2010-27.11.2011 produktów strukturyzowanych na polskim rynku kapitałowym (zob. rys. 1, tab. 1).



Rys. 1. Podział produktów według konstrukcji za okres 06.03.2011-06.03.2012

Źródło: Portal finansowy.

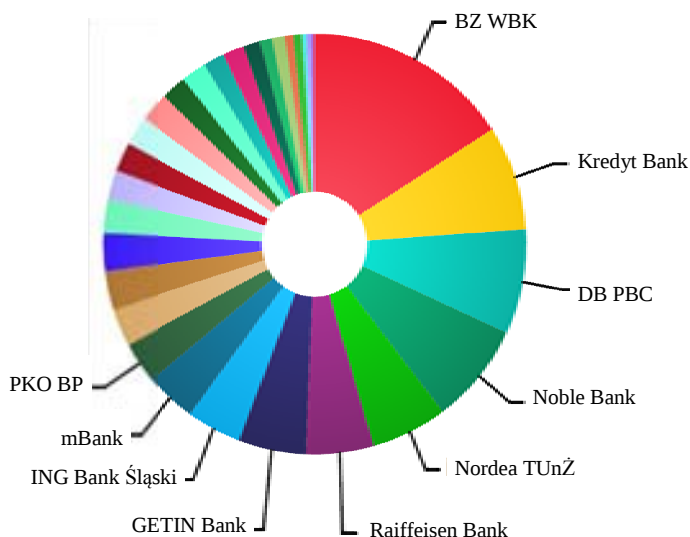
Tabela 1

Udział produktów strukturyzowanych w rynku, w zależności od przyjętej konstrukcji produktu strukturyzowanego w okresie 06.03.2011-06.03.2012

Konstrukcja produktu strukturyzowanego	Udział (%)
Polisa na życie z UFK	18,39
Polisa inwestycyjna	28,39
Lokata inwestycyjna	36,13
Certyfikat depozytowy	7,74
Obligacja strukturyzowana	3,23
BPW	2,58
FIZ	0,97
Fundusz zagraniczny	2,58

Źródło: Opracowanie na podstawie portalu finansowego.

Produkty strukturyzowane oferowane są przez liczne instytucje finansowe i banki. Największą ofertę w zakresie produktów strukturyzowanych w okresie 06.03.2011-06.03.2012 przedstawił Bank WBK, którego oferta stanowiła 15,9% wszystkich ofert na rynku w badanym okresie (rys. 2).



Rys. 2. Podział produktów według instytucji emitującej za okres 06.03.2011-06.03.2012

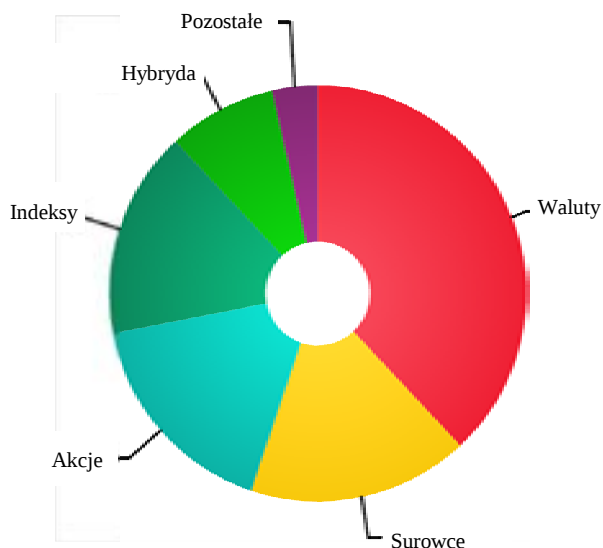
Źródło: Portal finansowy.

Produkty strukturyzowane zwykle charakteryzujemy ze względu na [7; 89; 10]:

1. Instrument bazowy, czyli instrument finansowy, na którym oparty jest produkt strukturyzowany. Mogą to być indeksy giełdowe (np. WIG 20 czy S&P 500), akcje wybranych spółek, ceny surowców, towarów rolnych, metali szlachetnych, walut obcych czy funduszy inwestycyjnych, a nawet pogoda. Bank, który tworzy produkt strukturyzowany wystawia opcję na instrument bazowy. Instrument ten w większości przypadków musi być więc konstruowany i prezentowany w sposób obiektywny, aby możliwa była wycena opcji (np. indeks notowany na giełdzie lub tworzony przez renomowaną instytucję finansową według określonych zasad)*.

W okresie 06.03.2011-06.03.2012 instrumentem bazowym produktów strukturyzowanych były najczęściej waluty (38,06%), co szczegółowo prezentuje rys. 3.

* <http://www.wealth.pl/akademia-lokat-strukturyzowanych/sloownik>, 2012.

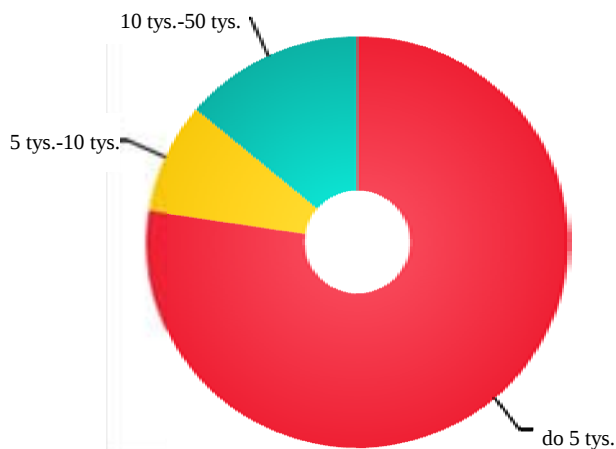


Rys. 3. Podział produktów według rodzajów aktywa za okres 06.03.2011-06.03.2012

Źródło: Portal finansowy.

2. Minimalną kwotę inwestycji.

W okresie 06.03.2011-06.03.2012 w 77,42% produkty były oferowane już przy inwestycji kapitału na poziomie 5000 zł (rys. 4).

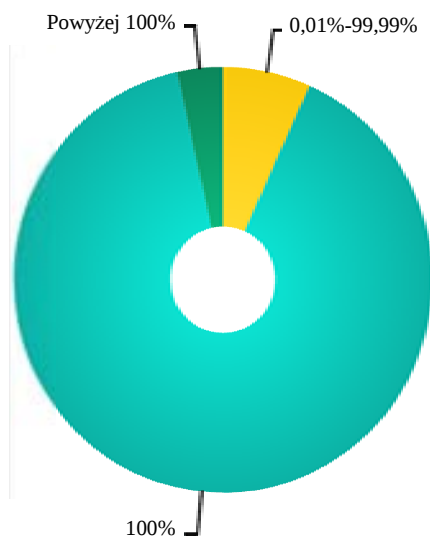


Rys. 4. Podział produktów według minimalnej wpłaty kapitału za okres 06.03.2011-06.03.2012

Źródło: Portal finansowy.

3. Ochronę zainwestowanego kapitału.

Jest to gwarancja zwrotu całości lub określonej części kwoty wpłaconej przez inwestora w ramach produktu strukturyzowanego. Najczęściej gwarantowany jest cały kapitał (gwarancja 100% kapitału). Zdarzają się jednak oferty z gwarancją na poziomie niższym od 100% oraz z wyższą (rys. 5).

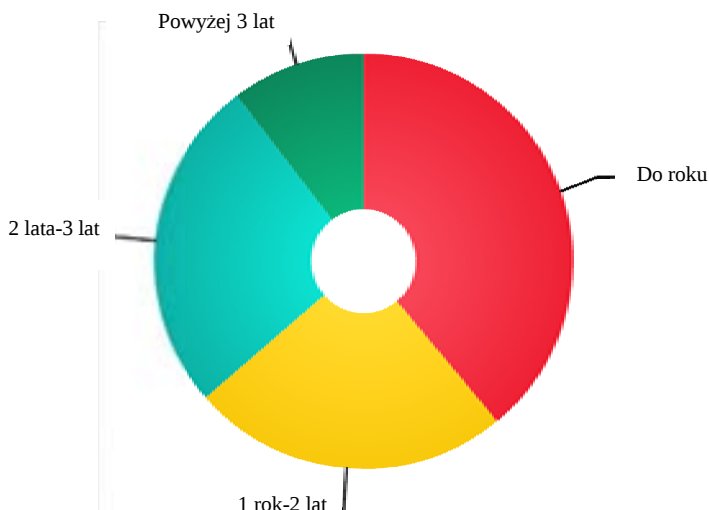


Rys. 5. Podział produktów według poziomu ochrony zainwestowanego kapitału za okres 06.03.2011-06.03.2012

Źródło: Portal finansowy.

4. Czas trwania.

Najczęściej produkty były oferowane na okres jednego roku (rys. 6).



Rys. 6. Podział produktów według czasu trwania za okres 06.03.2011-06.03.2012

Źródło: Portal finansowy.

5. Współczynnik partycypacji, czyli procentowy wskaźnik udziału inwestora w zmianie instrumentu bazowego (najczęściej w jego wzroście). W trakcie subskrypcji produktu strukturyzowanego partycypacja najczęściej podawana jest w przedziale np. od 100 do 130%, a jej ostateczna wartość podawana jest do wiadomości inwestorów po zakończeniu zapisów*.
6. Podatek (zob. rys. 7).

* <http://www.wealth.pl/akademia-lokat-strukturyzowanych/slownik>, 2012.



Rys. 7. Podział produktów według podatku za okres 06.03.2011-06.03.2012

Źródło: Portal finansowy.

1. Specyfikacja papieru wartościowego przyjętego do wyceny

BPW strukturyzowany przyjęty do wyceny jest strukturyzowanym, niezabezpieczonym papierem wartościowym na okaziciela, niemającym formy dokumentu (niematerialny). Wartość nominalna* oraz emisyjna** oferowanego BPW strukturyzowanego wynosi 1000 PLN za sztukę BPW, przy czym cena emisyjna ma sens w przypadku dojścia do skutku emisji BPW strukturyzowany, jedynie w przypadku osiągnięcia progu emisji. Informacja o niedojszcju emisji do skutku jest przekazana inwestorom:

- w ogłoszeniu zamieszczonym w dzienniku o zasięgu ogólnokrajowym,
- na stronie internetowej emitenta,
- na stronie internetowej odpowiedniego Domu Maklerskiego,
- w formie pisemnej.

Oferowany papier wartościowy jest inwestycją o rocznym horyzoncie czasowym, realizowaną zgodnie z datami przedstawionymi w tab. 2.

* „Wartość nominalna” oznacza kwotę stanowiącą iloczyn wartości nominalnej jednego BPW Strukturyzowanego wskazanej w warunkach emisji oraz liczby BPW strukturyzowanych zapisanych na rachunku BPW w dacie ustalenia praw z BPW.

** „Cena emisyjna” oznacza cenę za jeden BPW strukturyzowany wskazaną w warunkach emisji.

Tabela 2

Podstawy daty związane z inwestycją w Bankowy Papier Wartościowy Strukturyzowany

Okres subskrypcji*	Od 24.10.2010 r. do 18.11.2010
Data emisji	22 listopada 2011 r.,
Data wykupu**	22 listopada 2012 r.
Data Ustalenia kwoty wykupu	22 listopada 2012 r.

* Okres przyjmowania wpłat na produkt strukturyzowany. Po zakończeniu subskrypcji produkt jest konstruowany, za zebraną kwotę kupowane są obligacje oraz opcje (ewentualnie certyfikaty). W czasie subskrypcji wybrane parametry produktu nie są ściśle ustalone. Najczęściej dotyczy to wskaźnika partycypacji – podany jest zakres, w jakim ten wskaźnik się znajdzie. Ostateczna wartość partycypacji jest ustalana po zakończeniu zapisów, gdy znana już jest cena opcji, która zostanie kupiona. Subskrypcje trwają najczęściej od 2 do 5 tygodni.

** Z datą wykupu jest związana „cena odkupu”, która oznacza cenę za jeden BPW strukturyzowany, po której emitent odkupuje BPW strukturyzowane, wyznaczona zgodnie z zasadami określonymi w warunkach emisji.

W dacie wykupu emitent zapłaci posiadaczom bpw strukturyzowanych kwotę wykupu*. Do otrzymania kwoty wykupu są upoważnieni posiadacze, na których rachunku BPW w dacie ustalenia praw z BPW będą zapisane BPW strukturyzowane.

Jeżeli data płatności świadczenia pieniężnego z BPW przypada w dniu innym niż dzień roboczy, płatność świadczenia pieniężnego z BPW następuje pierwszego dnia roboczego po tym dniu, a posiadaczowi nie będzie przysługiwać roszczenie o odsetki za taki okres.

Premia z inwestycji będzie obliczona według następującej formuły:

$$P = L * PJ$$

gdzie:

P – premia,

L – liczba BPW strukturyzowanych zapisanych na rachunku BPW w dacie ustalenia praw z BPW,

PJ – premia jednostkowa, kwotowo od 1 sztuki BPW strukturyzowanego, zaokrąglona w dół do 1 złotego, obliczona według wzoru:

– w przypadku, gdy nie wystąpi osiągnięcie bariery:

$$PJ = N \cdot \text{MAX} \left[\frac{\text{Cena}_k}{\text{Cena}_p} - 1, 0 \right]$$

* „Kwota wykupu” oznacza sumę wartości nominalnej oraz premii płatną w dacie wykupu.

- w przypadku, gdy wystąpi osiągnięcie bariery, tzn., jeżeli w dniu obserwacji wskaźnika odniesienia, wskaźnik ten będzie równy lub większy od poziomu bariery:

$$PJ = N * Kupon$$

gdzie:

N – wartość nominalna 1 BPW strukturyzowanego,

$Cena_p$ – wartość wskaźnika odniesienia* w pierwszym dniu obserwacji wskaźnika odniesienia**, tj. 22.11.2011 r. (tab. 2),

$Cena_k$ – wartość wskaźnika odniesienia w ostatnim dniu obserwacji wskaźnika odniesienia, tj. 20.11.2012.

Ponadto:

$$Bariera = C_p \cdot x\%$$

$$Kupon = 10\%$$

gdzie $x\%$ wynosi 150-180%.

Dyspozycje zbycia BPW strukturyzowanych są rozliczane 2 dni robocze po dacie odkupu, tj., odkupione przez emitenta BPW strukturyzowane są wykسیgowane z rachunku BPW, a kwota odkupu jest przekazana na rachunek BPW lub na rachunek bankowy wskazany przez posiadacza zgodnie z umową o prowadzenie rachunku BPW. Minimalna liczba BPW strukturyzowanych będących przedmiotem jednej dyspozycji odkupu wynosi 5 sztuk. Cena odkupu jest ustalana w dacie odkupu, w oparciu o rynkową wycenę aktywów finansowych, w które lokowane są środki stanowiące wartość BPW strukturyzowanych.

Informacyjna cena odkupu jest zamieszczana na stronie internetowej emitenta i na stronie internetowej odpowiedniego Domu Maklerskiego. Podanie informacyjnej ceny odkupu nie zobowiązuje emitenta do odkupu ani posiadacza do sprzedaży BPW Strukturyzowanych po tej cenie. Posiadacz może złożyć dyspozycję zbycia BPW Strukturyzowanych:

- bez limitu ceny (odkup bezwarunkowy) – dyspozycja zbycia BPW strukturyzowanych jest realizowana przez emitenta po cenie odkupu, niezależnie od poziomu ceny odkupu, lub

* „Wskaźnik odniesienia” oznacza zmienną rynkową (np. stopę procentową, kurs walutowy, indeks rynku kapitałowego, cenę akcji, cenę towaru) wskazaną w warunkach emisji, od której, zgodnie z warunkami emisji danej emisji BPW strukturyzowanych uzależniona jest premia.

** W przypadku, gdy dany dzień obserwacji wskaźnika odniesienia nie jest potencjalnym dniem roboczym giełdy, dniem obserwacji wskaźnika odniesienia będzie najbliższy następny potencjalny dzień roboczy giełdy.

- z limitem ceny (odkup warunkowy) – dyspozycja zbycia BPW strukturyzowanych jest realizowana przez emitenta po cenie odkupu, o ile cena odkupu jest wyższa lub równa limitowi ceny określonego przez posiadacza w dyspozycji zbycia.

Emitent, za zgodą sądu, może złożyć do depozytu sądowego kwotę jakiegokolwiek świadczenia pieniężnego z BPW, jeżeli:

- na co najmniej 2 dni robocze przed datą wykupu lub datą wcześniejszego wykupu nie otrzymał od posiadacza informacji niezbędnych do dokonania płatności świadczenia pieniężnego z BPW, lub
- istnieje spór lub poważna wątpliwość co do tego, kto jest uprawniony z BPW strukturyzowanych, lub
- w jego uzasadnionej opinii dokonanie płatności świadczenia pieniężnego z BPW na rzecz posiadacza wiązałoby się z naruszeniem przepisów prawa lub naraziłoby emitenta na karę grzywny, inną karę bądź odpowiedzialność odszkodowawczą.

Złożenie kwot świadczenia pieniężnego z BPW do depozytu sądowego zwalnia emitenta z zobowiązań wobec posiadacza.

Wskaźnikiem odniesienia oferowanego przez emitenta BPW strukturyzowanego jest kurs wymiany EUR/ PLN, tzn. jeśli kurs wymiany EUR/PLN w dniu zapadalności inwestycji nie będzie równy lub wyższy od bariery znajdującej się na poziomie 150-180% wartości początkowej w dacie wykupu (22.11.2012 r.) klient otrzyma 100% partycypacji we wzroście ceny EUR/PLN na koniec okresu inwestycji. Jeśli w dniu zapadalności inwestycji kurs wymiany EUR/PLN będzie równy lub wyższy od bariery w wysokości 150-180% wartości początkowej, klient w dacie wykupu otrzyma zysk w wysokości 10% za okres inwestycji.

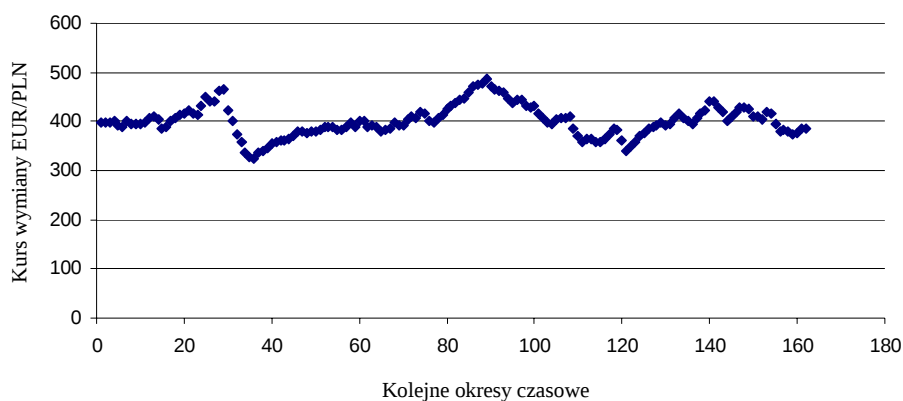
2. Oszacowanie inwestycji w BPW strukturyzowane

Celem oszacowania zysku z inwestycji przez inwestora, czy też przez emitenta należy przewidzieć przyszłą wartość wskaźnika, czyli w analizowanym produkcie kurs wymiany EUR/PLN, ponieważ przyszły zysk z inwestycji zależy ściśle od jego wartości w dniu zapadalności BPW.

Model wykorzystany do prognozowania wartości kursu EUR/PLN możemy opisać w kilku kluczowych etapach*:

ETAP 1

Podział szeregu prezentującego kurs wymiany EUR/PLN w okresie 01.1998-06.2011 na podszeregi n-elementowe.

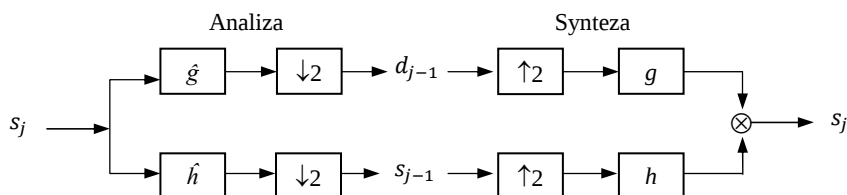


Rys. 8. Szereg prezentujący kurs EUR według NBP

Źródło: Dane GUS.

ETAP 2

Transformata falkowa podszeregów n-elementowych algorytmem „à trous”.

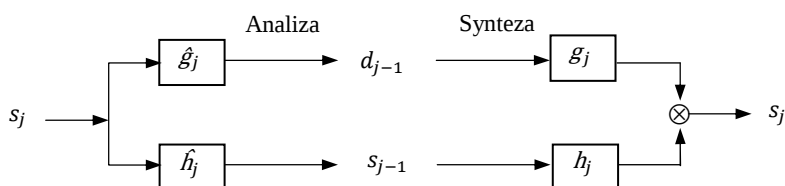


Rys. 9. Klasyczny schemat dekompozycji i rekonstrukcji falkowej (jeden etap)

* Wszelkie symulacje komputerowe i obliczenia w przedstawionym algorytmie wykonano w programie MATLAB korzystając z własnych programów autorskich.

Klasyczny schemat dekompozycji falkowej przedstawia rys. 10 (jeden etap analizy DWT i odpowiadający mu etap syntezy IDWT). Całościowy schemat może zawierać dowolną ilość etapów, ograniczeniem w przypadku transformacji blokowych jest długość bloku.

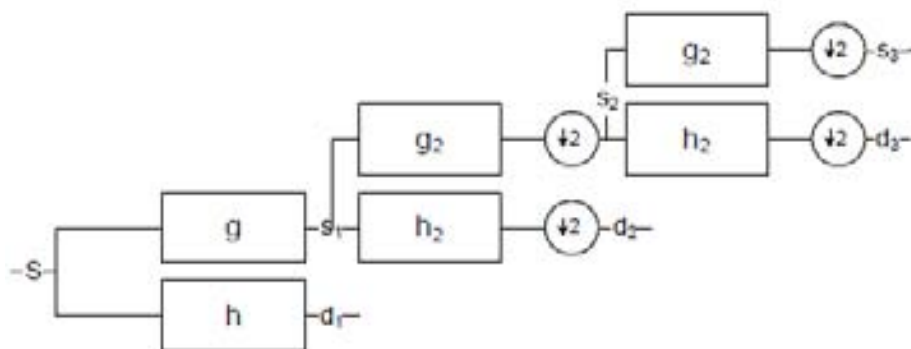
Klasyczny schemat dekompozycji falkowej uwzględnia diadyczność DWT (operacje podpróbkowania odpowiadają rozciąganiu funkcji falkowej i skalującej o czynnik 2). Diadyczność nie jest niestety automatyczna, gdy zrezygnujemy z decymacji, aby otrzymać wersję nadmiarową DWT, czyli reper. W takim przypadku musimy użyć zmodyfikowanych filtrów odpowiednio do etapu dekompozycji, tak aby uzyskać efekt rozciągania funkcji falkowej i skalującej. Najprostsze rozwiązanie daje tzw. algorytm „a trous”, w którym wprowadza się „nadpróbkowanie filtrów”, wstawiając zera pomiędzy współczynniki filtrów g_j i h_j .



Rys. 10. Niedecymowany schemat dekompozycji i rekonstrukcji falkowej (jeden etap)

Minimalną nadmiarowość uzyskuje się w wyniku użycia „filtrów nadpróbkowanych” (zgodnie z algorytmem „a trous”) i braku decymacji przynajmniej na pierwszym poziomie dekompozycji. Niedecymowana WT ma własności reperów w poszczególnych podprzestrzeniach detali z tą cechą charakterystyczną, że stanowi sumę mnogościową wielu baz Riesz, czyli liniowa zależność jego elementów jest cechą pochodną zwielokrotnienia liczebności o czynnik $2N$. Jeśli suma jego elementów daje element zerowy, to reper jest sztywny. W przypadku rozkładów ortogonalnych i biortogonalnych własność zerowania się sumy elementów repera jest „automatyczna”, jeśli przynajmniej na pierwszym etapie diadycznej dekompozycji falkowej nie dokonamy decymacji. Taka własność ma znaczenie w przypadku dekompozycji biortogonalnych, które w dziedzinie transformacji falkowej nie zachowują energii sygnału. Falki biortogonalne (zwłaszcza o wielu momentach znikających) w porównaniu do falek ortogonalnych uzyskują gorsze wyniki w dziedzinie odsumowania i kompresji stratnej. Własność zachowania energii jest ważna ze względu na to, że przy zastosowaniu progów lub kwantyzacji nie mamy jasno określonego kryterium i usunięcie relatywnie małego współczynnika falkowego może zaowocować dość dużą degradacją sygnału. Powyższa wada została częściowo skompensowana przez użycie algorytmu indeksowania ważnych współczynników falkowych.

Pod pojęciem „minimalnie nadpróbkowaną” przyjęto taką, która pomija decymację tylko na pierwszym etapie, a w kolejnych używa takich samych „dwukrotnie nadpróbkowanych” filtrów zgodnie z algorytmem „a trous” wraz z operacjami decymacji. W przypadku zastosowania schematu liftingowego w wersji „a trous” na każdym etapie dekompozycji filtry mają tę samą „dwukrotnie nadpróbkowaną” strukturę.



Rys. 11. Minimalnie nadpróbkowana DDWT

Pod pojęciem „wstępnie decymowanej” DDWT przyjęto taką, która wprowadza decymację tylko na kilku pierwszych etapach, a na kolejnych używa coraz wyżej „nadpróbkowanych” filtrów zgodnie z algorytmem „a trous”. Takie podejście prowadzi do transformacji, która jest mało wrażliwa na przesunięcie i umożliwia zastosowanie DDWT jako decymowanego wybielającego widmo filtra dopasowanego. Transformacja wykazuje wzrastającą wykładniczo nadmiarowość, począwszy od pierwszego etapu bez decymacji. Wstępna decymacja pozwala na znacząco mniejsze wykorzystanie zasobów FPGA (zwłaszcza pamięciowych) i stanowi kompromis pomiędzy rozdzielczością czasową detekcji a zajętością zasobów [1; 2; 3; 4; 5; 6; 9].

ETAP 3

Inicjalizacja sztucznej sieci neuronowej [11; 12; 13] na podstawie wygenerowanych na wcześniejszym etapie współczynników falkowych.

Najważniejszym elementem w neuronie są wartości poszczególnych wag, w których zapisana jest cała pamięć neuronu. Wagi mogą przyjmować wartości dodatnie, co reprezentuje synapsę pobudzającą, lub ujemne – reprezentujące synapsę hamującą. W przypadku braku połączenia między neuronami waga jest równa 0.

Wyróżniamy dwa podstawowe modele neuronu:

- liniowy,
- nieliniowy.

Neurony nieliniowe można podzielić dodatkowo na dwie podklasy:

- neuronów nieliniowych dyskretnych,
- neuronów nieliniowych ciągłych.

W obu przypadkach wartość sygnału na wyjściu neuronu obliczana jest na podstawie potencjału membranowego, czyli sumy ważonej sygnałów wejściowych i wag zgodnie ze wzorem:

$$y = f\left(\sum_{i=0}^N w_i x_i\right) = f(w^T x) = f(\text{net})$$

gdzie:

x – wektor sygnałów wejściowych,

w – wektor współczynników wagowych (**w₀ = -θ**),

f – funkcja aktywacji*,

y – sygnał wyjściowy.

Potencjał *net* skalowany jest poprzez funkcję aktywacji w jednostce zwanej blokiem aktywacji. W przypadku neuronu liniowego funkcja aktywacji jest liniowa i nie zmienia wartości uzyskanych z sumatora.

* W praktycznych zastosowaniach sieci neuronowych liniowa funkcja aktywacji może spowalniać działanie sieci, a uzyskanie precyzyjnych wyników jest bardzo mało prawdopodobne. Z tego powodu stosuje się nieliniowe funkcje aktywacji. Najczęściej stosowane nieliniowe funkcje aktywacji to: funkcja unipolarna, funkcja sigmoidalna, tangens hiperboliczny dane następującymi wzorami:

- funkcja unipolarna progowa McCullocha:

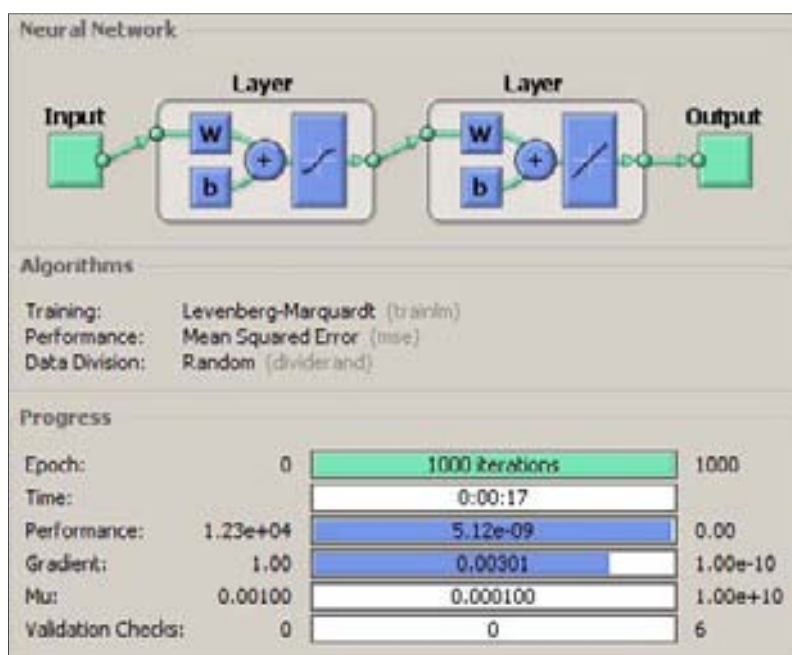
$$f(x) = \begin{cases} 1, & x \geq 0 \\ 0, & x < 0 \end{cases}$$

- Funkcja sigmoidalna:

$$f(x) = \frac{1}{1 + e^{-\beta x}}$$

- Tangens hiperboliczny:

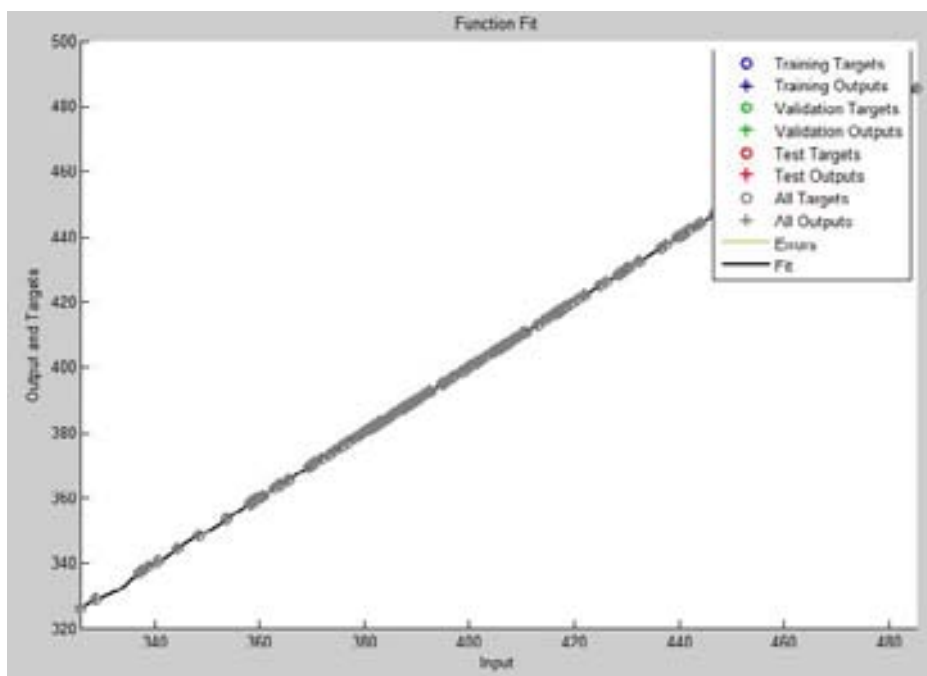
$$f(x) = \frac{2}{1 + e^{-\beta x}} - 1 = \frac{1 - e^{-\beta x}}{1 + e^{-\beta x}}$$



Rys. 12. Parametry uczenia sieci

Źródło: Opracowanie własne na podstawie badań.

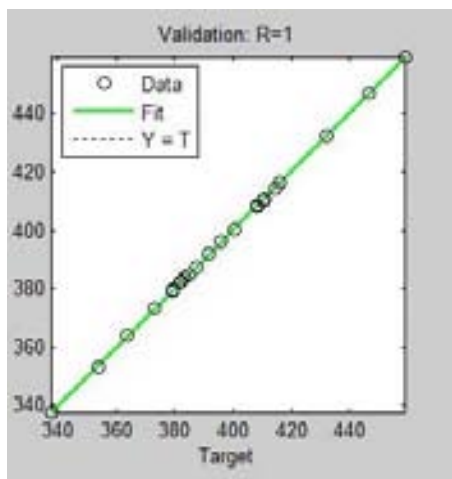
W rozważanym przypadku zastosowanie pojedynczego neuronu nie wystarczy do rozwiązania problemu, więc zbudowano sztuczną sieć neuronową (w skrócie SSN), określaną mianem systemu przetwarzania informacji symulującego działanie rzeczywistych, uczących się, funkcjonujących w mózgu struktur, której podstawowe parametry prezentuje rys. 12, a dopasowanie rys. 13.



Rys. 13. Dopasowanie SSN

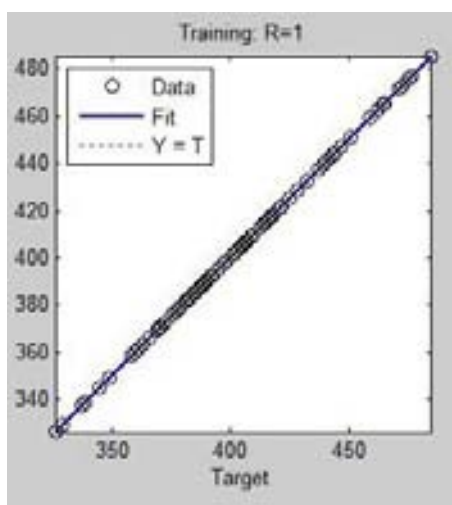
Źródło: Opracowanie własne na podstawie badań.

Efektywność uczenia poszczególnych zbiorów prezentują rys. 14, 15, 16, 17.



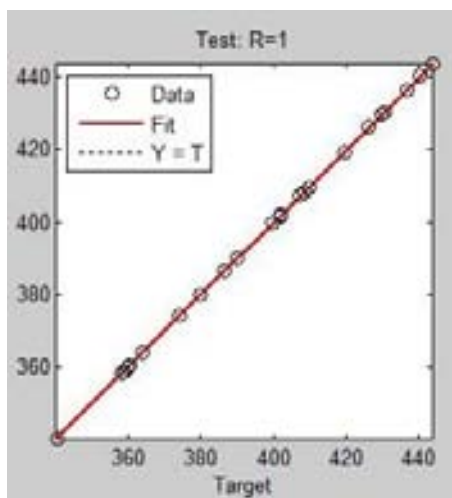
Rys. 14. Dopasowanie SSN

Źródło: Opracowanie własne na podstawie badań.



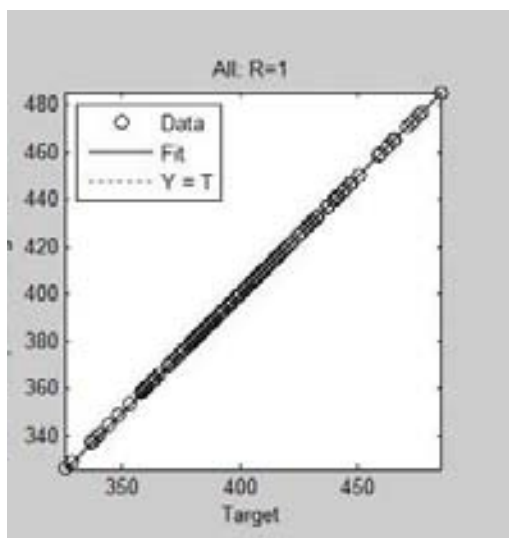
Rys. 15. Dopasowanie SSN

Źródło: Opracowanie własne na podstawie badań.



Rys. 16. Dopasowanie SSN

Źródło: Opracowanie własne na podstawie badań.



Rys. 17. Dopasowanie SSN

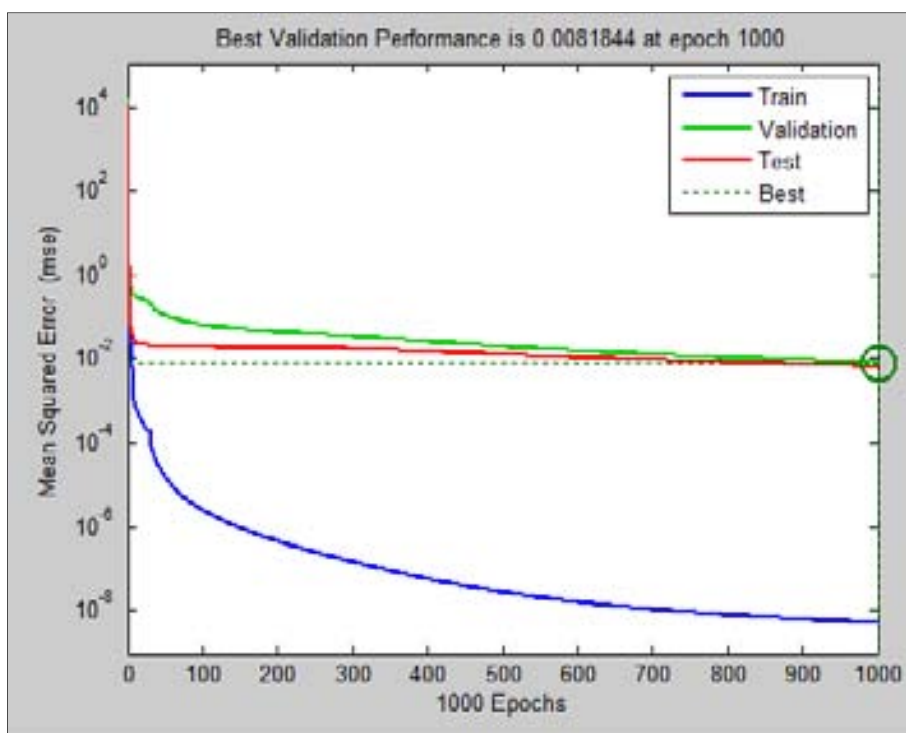
Źródło: Opracowanie własne na podstawie badań.

ETAP 4

Odwrotna transformata falkowa.

Efektem tego etapu jest uzyskanie prognozowanych 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 [1; 2; 3; 4; 5; 6; 9].

Wartość ta wynosi 4,79, jednak nie jest to wartość bezbłędna, ponieważ sztuczna sieć neuronowa na etapie uczenia oraz predykcji odpowiednich współczynników generuje błędy, odpowiednio: błąd uczenia sieci neuronowej i błąd testowy sztucznej sieci neuronowej, które przedstawia rys. 18.



Rys. 18. Błąd sztucznej sieci neuronowej dla generowanych współczynników falkowych

Źródło: Opracowanie własne na podstawie badań.

Dysponując przyszłą wartością kursu wymiany EUR/PLN inwestor może przewidzieć ewentualny zysk bądź stratę inwestycji. Mamy zatem:

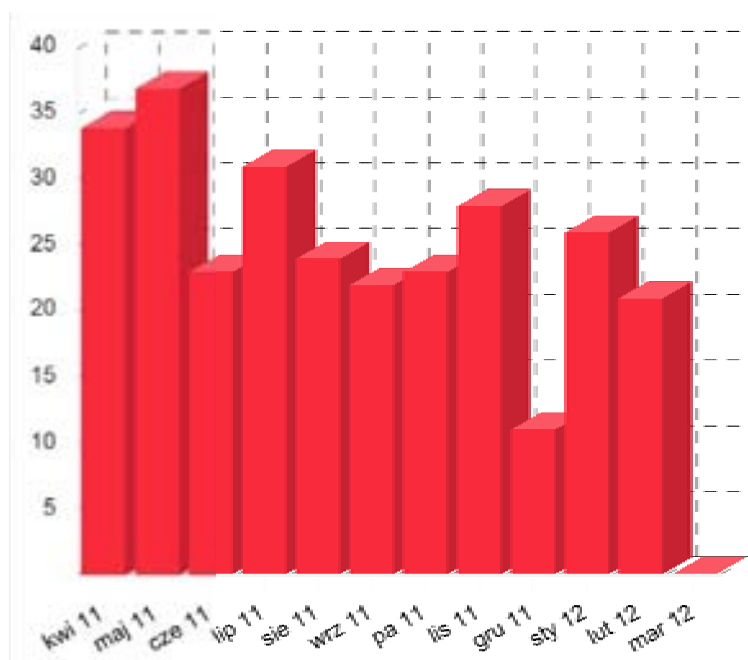
- przewidywana wartość wskaźnika na dzień 22.11.2012 r. – 4,79 PLN,
- wartość wskaźnika na dzień 22.11.2011 r. – 4,44 PLN,
- zmiana wskaźnika w stosunku do wartości początkowej: 107,88%.

Zgodnie z zasadami zaprezentowanymi w specyfikacji BPW strukturyzowanego, głoszącymi, że jeśli kurs wymiany EUR/PLN w dniu zapadalności inwestycji nie będzie równy lub wyższy od bariery znajdującej się na poziomie 150-180% wartości początkowej, w dacie wykupu (22.11.2012 r.) klient otrzyma 100% partycypacji we wzroście ceny EUR/PLN na koniec okresu inwestycji.

Inwestor otrzymuje zatem 7,88% zysku od zainwestowanego kapitału w BPW strukturyzowane pomniejszonego jedynie o opłatę dystrybucyjną.

Zakończenie

Przyszłość rynku produktów finansowych należy do produktów elastycznych, umożliwiających znalezienie złotego środka pomiędzy możliwościami osiągnięcia satysfakcjonujących zysków, a jednocześnie chroniących kapitał. Produkty strukturyzowane doskonale spełniają tę rolę. Doceniają to klienci, którzy mogą zarabiać bez względu na panującą na rynku koniunkturę, co prezentuje rys. 19.



Rys. 19. Liczba produktów strukturyzowanych oferowanych w okresie 04.2011-02.2012

W przedstawionym oszacowaniu BPW przyjęto opcję, że podczas trwania produktu nie wystąpią zakłócenia wskaźnika. Gdyby jednak w dniu obserwacji wskaźnika odniesienia wystąpiło w stosunku do indeksu przynajmniej jedno z zakłóceń indeksu, wówczas emitent musiałby podjąć w odniesieniu do BPW strukturyzowanych takie czynności zastępcze, jakie w odniesieniu do transakcji zabezpieczającej podjąłby podmiot zabezpieczający, przy czym podmiot zabezpieczający mógłby nie podjąć żadnej z czynności zastępczych. W przypadku, gdy dla danej emisji BPW strukturyzowanych występuje kilka podmiotów zabezpieczających i podmioty te podejmą różne czynności zastępcze lub część z nich nie podejmie czynności zastępczych po wystąpieniu danego zakłócenia

indeksu, emitent stosuje te czynności zastępcze w proporcjach, w jakich zawarł z podmiotami zabezpieczającymi transakcje zabezpieczające dla danej emisji BPW strukturyzowanych.

Literatura

1. Aussum.et al. Combing neural network forecast on wavelet-transformed series [J]. Connection Science, 1997.
2. Burrus C.S., Gopinath R.A., Guo H.: *Introduction to Wavelets and Wavelet Transform*, Prentice Hall 1998.
3. Chui Ch.K., *Wavelets: A Mathematical Tool for Signal Analysis*, SIAM 1997.
4. Dyduch M., *Zastosowanie sieci falkowo - neuronowej do predykcji ekonomicznych szeregów czasowych*, w: *Prognozowanie w zarządzaniu firmą*, red. P. Dittmann, J. Krupowicz, Akademia Ekonomiczna, Wrocław 2006.
5. Dyduch M., *Sieć falkowo-neuronowa w środowisku ekonomicznym*, w: *Zarządzanie – Finanse – Ekonomia*, Warsztaty doktorskie '06, red. T. Trzaskalik, Akademia Ekonomiczna, Katowice 2007.
6. Dyduch M., *Sieć falkowo-neuronowa jako skuteczne narzędzie do analizy i predykcji szeregów czasowych*, w: *Metody matematyczne, ekonometryczne i komputerowe w finansach i ubezpieczeniach 2006*, red. P. Chrzan, T. Czernik, Akademia Ekonomiczna, Katowice 2008.
7. Dyduch M., *Przestrzenno-czasowa analiza atrakcyjności finansowych produktów strukturyzowanych*, *Ekonometria Przestrzenna i Regionalne Analizy Ekonomiczne*, Łódź 2010.
8. Dyduch M., *Znaczenie produktów strukturyzowanych w Polsce*, *Ekonomia. Finanse. Współczesne wyzwania i kierunki rozwoju*, Uniwersytet Ekonomiczny, Katowice 2010.
9. Dyduch M., *Współczynniki transformaty falkowej jako narzędzie generujące prognozę przedziałową szeregów czasowych*, w: *Modelowanie preferencji a ryzyko'10*, red. T. Trzaskalik, Akademia Ekonomiczna, Katowice 2010.
10. Dyduch M., *Analiza porównawcza produktów strukturyzowanych wyemitowanych na polskim rynku finansowym w latach 2000-2010*, Materiały krakowskiej konferencji młodych uczonych 2011, Sympozja i Konferencje KKMU nr 6, Fundacja Studentów dla AGH, Grupa Naukowa Pro Futro, Kraków 2011.
11. Rutkowska D., Piliński M., Rutkowski M., *Sieci neuronowe, algorytmy genetyczne i systemy rozmyte*, Wydawnictwo Naukowe PWN, Warszawa 1997.
12. Tadeusiewicz R., *Sieci neuronowe*, Akademicka Oficyna Wydawnicza RM, Warszawa 1993.
13. Witkowska D., *Sztuczne sieci neuronowe i metody statystyczne. Wybrane zagadnienia finansowe*, Wydawnictwo C.H. Beck, Warszawa 2002.

BANKING STRUCTURED SECURITIES

Summary

The article presents a short description of BPW, which is a modern financial instrument, opening in front of customers to invest with full capital protection on both traditional and alternative markets (eg, stocks, commodities, currencies, real estate, commodities). BPW are a combination of deposit, to improve return on capital at the date of maturity of the investment and, working on potential profits. Part of the investment is for the purchase of financial instruments (usually an option), which will allow for the implementation of a specific investment strategy described in detail in terms of the issue of the BPW. In addition, BPW is subject to guarantees of the Bank Guarantee Fund.

Iwona Dittmann

Uniwersytet Ekonomiczny we Wrocławiu

PODOBIENSTWO ZMIAN ŚREDNICH CEN TRANSAKCYJNYCH 1 m² POWIERZCHNI MIESZKAŃ W WYBRANYCH MIASTACH WOJEWÓDZTWA ŚLĄSKIEGO

Wprowadzenie

W działalności podmiotów działających na rynku nieruchomości mieszkaniowych pomocne mogą być analizy kształtowania się cen transakcyjnych na rynkach lokalnych. Zagadnienie to poruszali w swoich pracach m.in. Żelazowski [14], Nykiel [7], Trojanek [12].

Na podstawie wcześniej przeprowadzonych badań dotyczących kształtowania się średnich cen transakcyjnych 1 m² powierzchni mieszkań na największych rynkach lokalnych w Polsce (warszawskim, katowickim, łódzkim, wrocławskim, gdańskim, białostockim, krakowskim i poznańskim) stwierdzono m.in., że występuje podobieństwo w zachodzeniu zmian cen w czasie na poszczególnych rynkach. W szczególności zaobserwowano, że zmiany średnich cen transakcyjnych na niektórych rynkach lokalnych występują równocześnie ze zmianami na innych rynkach lokalnych, naśladując z różnym opóźnieniem zmiany cen zachodzące na innych rynkach lokalnych bądź wyprzedzają zmiany zachodzące na innych rynkach lokalnych [4]. Postanowiono zatem zbadać czy podobne zjawiska zachodzą na mniejszych rynkach nieruchomości zlokalizowanych w jednym województwie. Badanie, którego wyniki zamieszczono w niniejszym opracowaniu dotyczyło kształtowania się średnich cen transakcyjnych 1 m² powierzchni mieszkań w wybranych (większych) miastach województwa śląskiego: Katowicach, Będzinie, Bielsku-Białej, Częstochowie, Jastrzębie-Zdroju, Jaworznie, Mikołowie, Piekarach Śląskich, Raciborzu, Ryb-

niku, Tarnowskich Górach, Tychach, Zawierciu i Żorach w okresie I kw. 2006 r.-II kw. 2011 r. Źródłem danych o kształtowaniu się cen był raport AMRON.

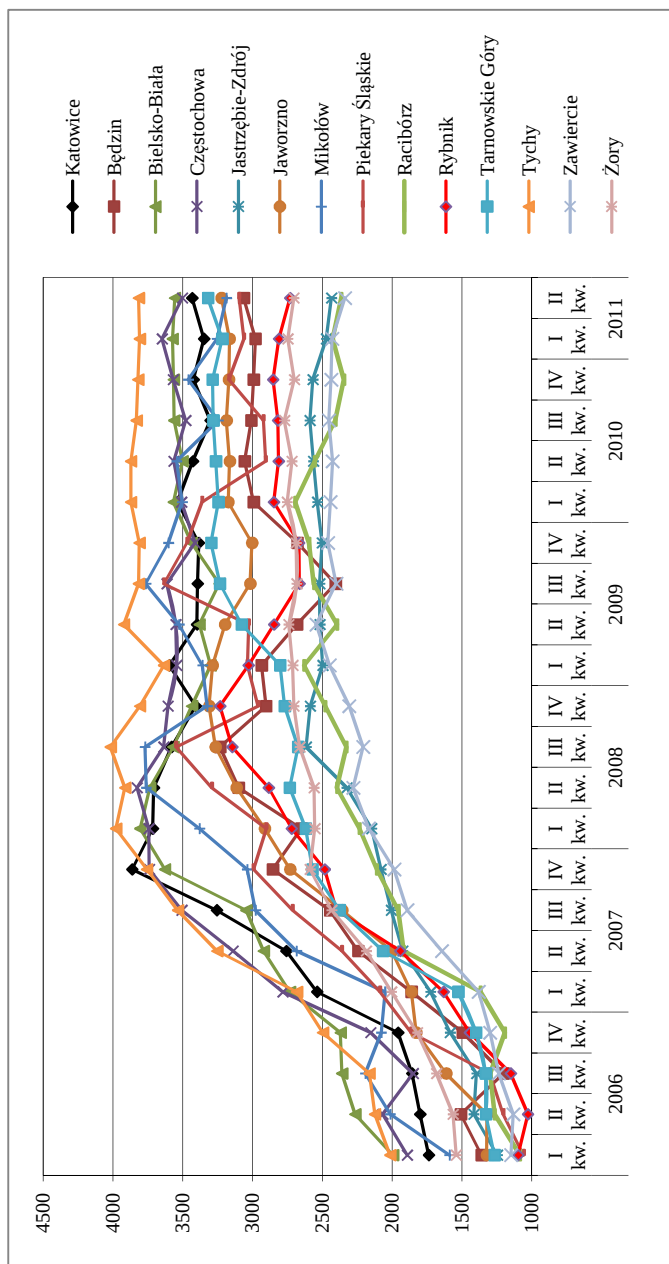
W badaniu postawiono hipotezy, że na analizowanych lokalnych rynkach nieruchomości w okresie objętym badaniem:

- 1) nie występowało podobieństwo poziomu średnich cen transakcyjnych 1 m² powierzchni mieszkań,
- 2) występowało zjawisko sigma konwergencji,
- 3) występowało podobieństwo w zachodzeniu zmian średnich cen transakcyjnych 1 m² powierzchni mieszkań w czasie (podobieństwo kształtu).

1. Podobieństwo poziomu średnich cen transakcyjnych 1 m² powierzchni mieszkań

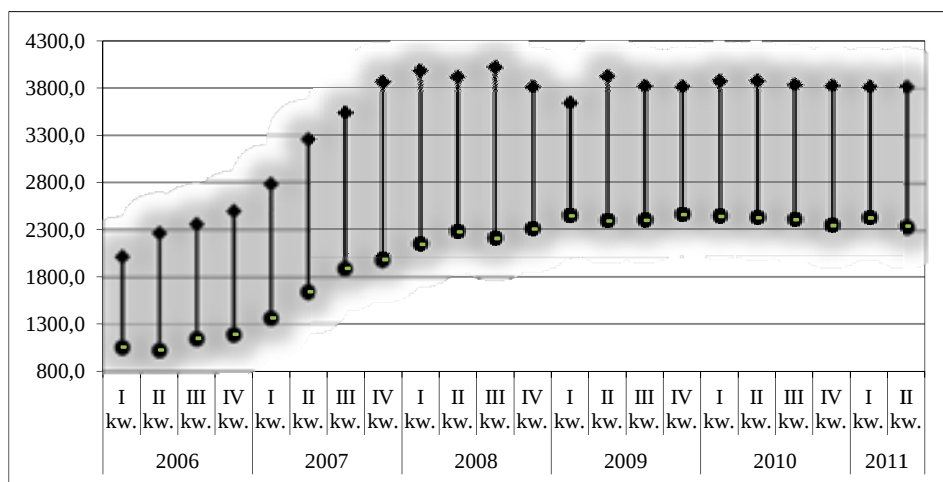
Podobieństwo poziomu średnich cen transakcyjnych 1 m² powierzchni mieszkań będzie występowało wówczas, gdy w każdym okresie różnice między maksymalnymi a minimalnymi cenami będą mniejsze od przyjętej wielkości progowej h^* . Wartości cen transakcyjnych będą się zatem mieściły w korytarzu o rozpiętości pionowej nie większej od h^* .

Kształtowanie się średnich cen transakcyjnych 1 m² powierzchni mieszkań na badanych rynkach mieszkaniowych przedstawiono na rys. 1. Wydaje się, że na jego podstawie można stwierdzić, że w badanym okresie występował brak podobieństwa poziomu średnich cen transakcyjnych 1 m² powierzchni mieszkań. W celu sprawdzenia tego przypuszczenia sporządzono rys. 2, na którym zamieszczono minimalne i maksymalne wartości średnich cen transakcyjnych w poszczególnych kwartałach oraz rys. 3, na którym zamieszczono obliczone rozstępy (różnice między maksymalnymi a minimalnymi cenami transakcyjnymi) dla każdego kwartału. Uwzględniając największy rozstępy (1900 zł) należałoby przyjąć wartość progową h^* powyżej 1900 zł, aby uznać, że występowało podobieństwo poziomu średnich cen transakcyjnych 1 m² powierzchni mieszkaniowej. Wydaje się jednak, że w porównaniu z obserwowanymi w badanym okresie cenami – nawet najwyższymi – jest ona w stosunku do nich zbyt wysoka, aby przyjąć ją za h^* . Należy zatem stwierdzić brak podobieństwa poziomu w analizowanej grupie miast.

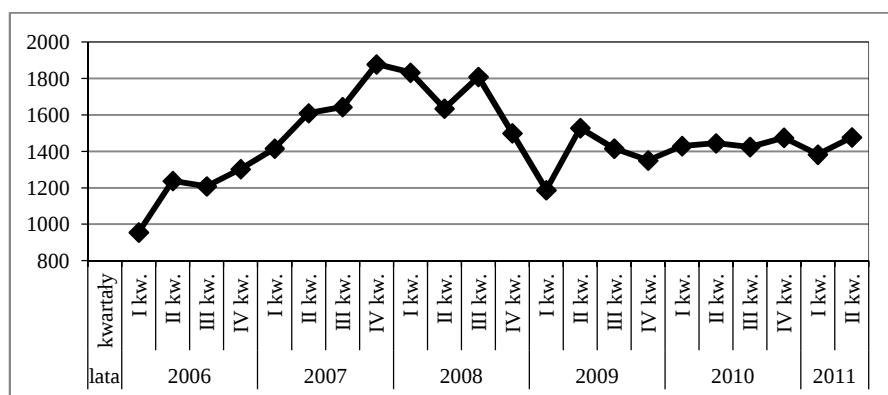


Rys. 1. Średnie ceny transakcyjne 1 m² powierzchni mieszkań (I kw. 2006 r.-II kw. 2011 r.) w zł

Źródło: Opracowanie własne na podstawie [9].



Rys. 2. Minimalne i maksymalne średnie ceny transakcyjne 1 m² powierzchni mieszkań (I kw. 2006 r.-II kw. 2011 r.) w zł



Rys. 3. Rozstępy średnich cen transakcyjnych 1 m² powierzchni mieszkań (I kw. 2006 r.-II kw. 2011 r.) w zł

2. Sigma konwergencja lokalnych rynków nieruchomości

Pojęcie „konwergencji” pojawia się na gruncie makroekonomii i najczęściej łączy się z procesem wzrostu gospodarczego i wyrównywaniem różnic między obiektami (krajami lub regionami) o różnym poziomie rozwoju ekonomicznego [1; 8; 10; 13]. Problem konwergencji, czyli upodabniania się badanych procesów (wyrównywania różnic) można rozważać w dwóch ujęciach, określanych mianem sigma konwergencji i beta konwergencji.

Sigma konwergencja występuje wówczas, gdy zróżnicowanie badanego zjawiska zmniejsza się w czasie, tzn. zachodzi relacja:

$$\sigma_t > \sigma_{t+T} \quad \text{lub} \quad V_t > V_{t+T} \quad (1)$$

gdzie:

$\sigma_t, (V_t)$ – odchylenie standardowe (współczynnik zmienności) miernika badanego procesu (np. PKB *per capita*) w okresie/momencie t ,

$\sigma_{t+T}, (V_{t+T})$ – odchylenie standardowe (współczynnik zmienności) miernika badanego procesu (np. PKB *per capita*) w okresie/momencie $t+T$,

T – długość rozważanego okresu.

W przypadku wystąpienia nierówności o kierunku przeciwnym niż we wzorze (1) mówimy, że występuje sigma dywergencja. Brak wyraźnych różnic pomiędzy odchyleniami standardowymi lub współczynnikami zmienności wskazuje na ustabilizowanie się procesu w czasie (w analizowanym okresie nie występuje ani sigma konwergencja, ani sigma dywergencja).

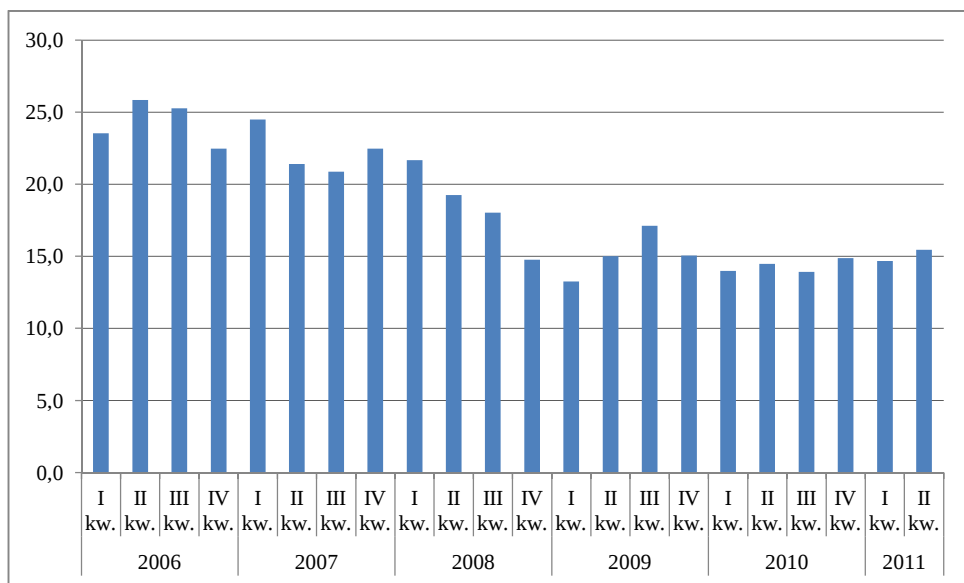
Zjawisko beta konwergencji występuje wówczas, gdy obiekty o niższej wartości początkowej miernika analizowanego procesu charakteryzują się wyższym tempem jego wzrostu niż podmioty o wyższej wartości początkowej tego miernika. „Maruder” (obiekt o niższej wartości początkowej miernika) dogania lidera upodabniając się do niego. Jeśli natomiast występuje zjawisko odwrotne, tzn. obiekty o niższej wartości początkowej miernika analizowanego procesu charakteryzują się niższym tempem jego wzrostu niż obiekty o wyższej wartości początkowej tego miernika („maruder” pozostaje coraz bardziej w tyle), występuje zjawisko beta dywergencji.

Obydwie koncepcje konwergencji są ze sobą powiązane. Sigma konwergencję możemy rozumieć jako zmianę zróżnicowania poziomu analizowanego miernika w czasie wokół średniej. Beta konwergencja informuje czy uporządkowanie jednostek według poziomu miernika zmienia się w czasie. Wystąpienie beta konwergencji jest warunkiem koniecznym, ale niedostatecznym dla wystąpienia sigma konwergencji. Oba podejścia są stosowane w analizie konwergencji i dywergencji także innych procesów, np. demograficznych [6; 11] czy społecznych [2]. Wydaje się, że mogą być również stosowane w analizach procesów zachodzących na rynkach nieruchomości.

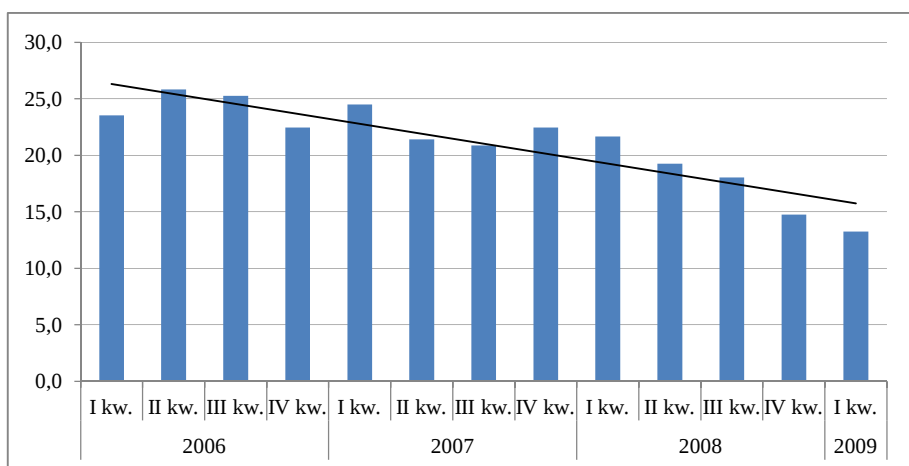
W celu zbadania czy występuje sigma konwergencja analizowanych miast województwa śląskiego pod względem średnich cen transakcyjnych 1 m² powierzchni mieszkań obliczono współczynniki zmienności cen. Zdecydowano się na użycie współczynników zmienności, a nie odchyleń standardowych. Na pod-

stawie wcześniejszych analiz stwierdzono, że współczynnik zmienności i odchylenie standardowe mogą dawać przeciwne wskazania odnośnie do występowania konwergencji lub dywergencji na analizowanych rynkach, zwłaszcza w przypadku silnych wzrostów lub spadków cen transakcyjnych [5]. Wydaje się, że większe znaczenie powinno się przywiązywać w tym przypadku do wskazań względnego miernika różnicowania (współczynnika zmienności) niż do miernika bezwzględnego (odchylenia standardowego).

Na podstawie wykresu (rys. 4) stwierdzono, że w kształtowaniu się w analizowanym okresie wartości współczynników zmienności średnich cen transakcyjnych 1 m² powierzchni mieszkań można wyróżnić dwa podokresy. W pierwszym (I kw. 2006 r.-I kw. 2009 r.) różnicowanie średnich cen transakcyjnych 1 m² powierzchni mieszkań maleje (występuje konwergencja), w drugim (II kw. 2009 r.-II kw. 2011 r.) występuje stabilizacja różnicowania (nie występuje ani konwergencja, ani dywergencja). Dla potwierdzenia tego przypuszczenia sporządzono rys. 5 i 6 oraz oszacowano parametry liniowych funkcji trendu (tab. 1 i 2).



Rys. 4. Współczynniki zmienności średnich cen transakcyjnych 1 m² powierzchni mieszkań dla poszczególnych kwartałów w %

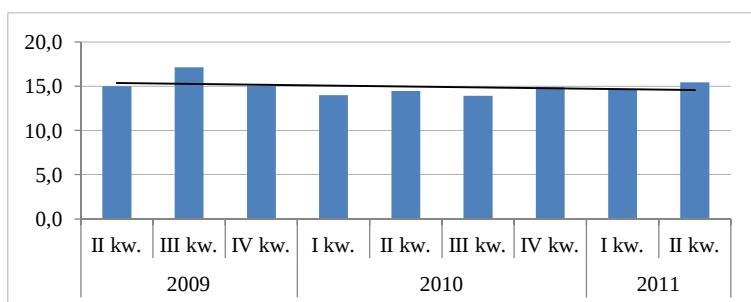


Rys. 5. Współczynniki zmienności średnich cen transakcyjnych 1 m² powierzchni mieszkań (I kw. 2006 r.-I kw. 2009 r.) w % wraz z funkcją trendu

Tabela 1

Wyniki estymacji funkcji trendu (I kw. 2006 r.-I kw. 2009 r.)

Zmienna zależna: współczynnik zmienności ceny 1m ² powierzchni mieszkaniowej				
R = 0,895		R ² = 0,801		
Standardowy błąd estymacji = 1,782				
Parametry		błąd standardowy	t stat	wartość-p
Przecięcie	27,182	1,048	25,924	3,255E-11
Współczynnik kierunkowy	-0,880	0,132	-6,662	3,552E-05



Rys. 6. Współczynniki zmienności średnich cen transakcyjnych 1 m² powierzchni mieszkań (II kw. 2009 r.-II kw. 2011 r.) w % wraz z funkcją trendu

Tabela 2

Wyniki estymacji funkcji trendu (II kw. 2009 r.-II kw. 2011 r.)

Zmienna zależna: współczynnik zmienności ceny 1 m ² powierzchni mieszkaniowej				
R = 0,284		R ² = 0,080		
Standardowy błąd estymacji = 0,979				
Parametry		błąd standardowy	t stat	wartość-p
Przecięcie	15,445	0,711	21,712	1,10895E-07
Współczynnik kierunkowy	−0,099	0,126	−0,785	0,458188156

Statystycznie istotnie różniący się od zera (wartość-p = 3,552E-05) ujemny współczynnik kierunkowy funkcji trendu (-0,880) i dobre dopasowanie modelu ($R^2 = 0,801$) w pierwszym podokresie oraz nieistotnie różniący się od zera (wartość-p = 0,458188156) współczynnik kierunkowy funkcji trendu (-0,099) i słabe dopasowanie modelu ($R^2 = 0,080$) potwierdziły wcześniejsze przypuszczenia o występowaniu konwergencji lokalnych rynków nieruchomości w okresie I kw. 2006 r.-I kw. 2009 r. oraz o stabilizacji w okresie II kw. 2009 r.-II kw. 2011 r.

3. Podobieństwo w zachodzeniu zmian średnich cen transakcyjnych 1 m² powierzchni mieszkań

W kształtowaniu się średnich cen transakcyjnych 1m² powierzchni mieszkań w analizowanych miastach można zauważyć podobne tendencje (rys. 1). Można zatem uznać, że występuje tzw. podobieństwo kształtu. Do jego pomiaru zastosowano miarę podobieństwa funkcji [3]. Otrzymane wyniki zamieszczono w tab. 3.

Za progową wartość miary podobieństwa przyjęto $m^* = 0,5$. Na tej podstawie za podobne pod względem zmian cen 1 m^2 uznano następujące rynki lokalne:

- Katowice i Piekary Śląskie,
- Częstochowa i Jastrzębie-Zdrój,
- Jaworzno i Rybnik,
- Zawiercie i Rybnik,
- Zawiercie i Tarnowskie Góry,
- Zawiercie i Żory.

Wśród największych lokalnych rynków mieszkaniowych w Polsce (8 miast wojewódzkich) stwierdzono występowanie:

- rynków zbieżnych (tzn. takich, na których zmiany cen transakcyjnych zachodziły równocześnie ze zmianami cen na innych rynkach),
- rynków wiodących (tzn. takich, na których zmiany cen transakcyjnych zachodziły wcześniej niż zmiany cen na innych rynkach),
- rynków naśladowujących (tzn. takich, na których zmiany cen transakcyjnych zachodziły później niż zmiany cen na innych rynkach) [4; 5].

Uzyskane w badaniach wyniki zamieszczono w tab. 4. W nawiasach podano wartości miar podobieństwa dla przedziałów o długości 22, 21, 20, 19 obserwacji (w przypadku rynków zbieżnych), dla przedziałów o długości 21, 20, 19 obserwacji (w przypadku rynków „przesuniętych” o 1 kwartał), dla przedziałów o długości 20, 19 obserwacji (w przypadku rynków „przesuniętych” o 2 kwartały).

Tabela 4

Wybrane miasta wojewódzkie – rynki zbieżne, wyprzedzające i naśladowujące

Rynki	Rynki				
	naśladowujące z opóźnieniem o 2 kwartały	naśladowujące z opóźnieniem o 1 kwartał	zbieżne	wyprzedzające o 1 kwartał	wyprzedzające o 2 kwartały
1	2	3	4	5	6
Warszawa	–	Kraków (0,671; 0,654; 0,635)	Łódź (0,705; 0,696; 0,681; 0,663)	Wrocław (0,694; 0,678; 0,660)	–
Wrocław	Kraków (0,787; 0,775)	Warszawa (0,694; 0,678; 0,66)	Gdańsk (0,615; 0,595; 0,574; 0,551) Poznań (0,515; 0,59; 0,569; 0,546)	–	Katowice (0,674; 0,657)

cd. tabeli 4

1	2	3	4	5	6
Białystok	Łódź (0,535; 0,51)	Kraków (0,747; 0,734; 0,719)	–	Gdańsk (0,764; 0,751; 0,738) Wrocław (0,565; 0,542; 0,516)	Poznań (0,667; 0,759) Katowice (0,661; 0,643)
Poznań	Białystok (0,667; 0,759)	–	Katowice (0,518; 0,593; 0,572; 0,549) Wrocław (0,515; 0,59; 0,569; 0,546)	–	–
Katowice	Wrocław (0,674; 0,657) Gdańsk (0,663; 0,644) Białystok (0,661; 0,643) Kraków (0,549; 0,521)	–	Poznań (0,518; 0,593; 0,572; 0,549)	–	–
Gdańsk	Kraków (0,696; 0,679)	Białystok (0,764; 0,751; 0,738)	Wrocław (0,615; 0,595; 0,574; 0,595)	–	Katowice (0,663; 0,644)
Kraków	–	–	Łódź (0,682; 0,672; 0,655; 0,636)	Białystok (0,747; 0,734; 0,719) Warszawa (0,671; 0,654; 0,635)	Wrocław (0,787; 0,775) Gdańsk (0,696; 0,679) Poznań (0,549; 0,634) Katowice (0,549; 0,521)
Łódź	–	–	Warszawa (0,705; 0,696; 0,681; 0,663) Kraków (0,682; 0,672; 0,655; 0,636)	Wrocław (0,595; 0,574; 0,55)	Białystok (0,535; 0,51)

Źródło: [5].

Postanowiono zatem sprawdzić, czy prawidłowości tego rodzaju występują także wśród badanych miast województwa śląskiego. W tym celu obliczono wartości miar podobieństwa funkcji dla danych opóźnionych odpowiednio o: 1 kwartał, 2 kwartały, 3 kwartały, 4 kwartały i 5 kwartałów. Okazało się, że przy przyjętej wartości progowej $m^* = 0,5$ występują takie podobieństwa między miastami. Wyniki przedstawiono w tab. 5. W przypadku stwierdzenia podobieństwa dla różnych przesunięć w czasie wybierano te o największych wartościach miary podobieństwa. Liczby w nawiasach są wartościami miar podobieństwa uszeregowanymi od przedziału złożonego z największej możliwej liczby obserwacji (najdłuższego) do przedziału złożonego z najmniejszej przyjętej liczby obserwacji (najkrótszego). Przykładowo w przypadku braku opó-

nienia wartości miar odnoszą się kolejno do przedziałów złożonych z: 22, 21, 20, 19, 18, 17 obserwacji; w przypadku opóźnienia o 1 kwartał odnoszą się do przedziałów złożonych z: 21, 20, 19, 18, 17 obserwacji. W pewnym stopniu ukazują one stabilność podobieństwa w czasie. W przypadku, gdy wartość miary niewiele przekraczała 0,5, w pojedynczych przypadkach na ogół jednak nie osiągała progu 0,5, pomijano te pary rynków lokalnych i nie zamieszczano ich w tabeli.

Na podstawie wyników zamieszczonych w tab. 5 można stwierdzić m.in. że:

- Katowice (miasto wojewódzkie) nie są rynkiem naśladowującym w stosunku do żadnego z badanych miast powiatowych, przeciwnie – stanowią rynek wyprzedzający w stosunku do Piekar Śląskich* (wyprzedzenie o 2 kwartały), Jaworzna i Raciborza (wyprzedzenie o 4 kwartały),
- podobieństwo między badanymi rynkami lokalnymi zachodzi nie tylko w tym samym czasie, ale także w przypadku małych i większych przesunięć w czasie; największe analizowane w pracy opóźnienia wynosiły 5 kwartałów, nie badano podobieństwa rynków przy większym opóźnieniu,
- w przypadku każdego rynku lokalnego można zidentyfikować rynki podobne, ich liczba waha się od 5 (Katowice, Będzin) do 11 (Piekary Śląskie, Rybnik),
- występują miasta, których żadne z badanych miast nie wyprzedza (Katowice) oraz miasta, których żadne z badanych miast nie naśladuje (Mikołów),
- wstępnie można zaobserwować, że liczba „par” i wartości miar podobieństwa dla rynków podobnych nie dają silnych podstaw do postawienia hipotezy o zachodzącym większym podobieństwie zmiany cen na największych rynkach lokalnych w danym regionie niż podobieństwie zachodzącym pomiędzy największymi rynkami miast wojewódzkich.

* Wartość miary podobieństwa była wyższa niż w przypadku braku opóźnienia.

Tabela 5

Wybrane miasta w województwie śląskim – rynki zbieżne, wyprzedzające i naśladowujące

Miasta naśladowujące	Opóźnienie	Miasta	Miasta naśladowujące	Opóźnienie	Miasta
1	2	3	4	5	6
Katowice	bez opóźnienia	–	Piekary Śląskie	bez opóźnienia	–
	o 1 kwartał	–		o 1 kwartał	–
	o 2 kwartały	–		o 2 kwartały	Katowice (0,557; 0,533; 0,506; 0,476), Tychy (0,523; 0,497; 0,468; 0,435)
	o 3 kwartały	–		o 3 kwartały	Rybnik (0,542; 0,633; 0,610), Żory (0,543; 0,518; 0,488)
	o 4 kwartały	–		o 4 kwartały	–
	o 5 kwartałów	–		o 5 kwartały	Częstochowa (0,500)
Będzin	bez opóźnienia	–	Racibórz	bez opóźnienia	–
	o 1 kwartał	Racibórz (0,592; 0,676; 0,659; 0,756; 0,741)		o 1 kwartał	–
	o 2 kwartały	–		o 2 kwartały	Piekary Śląskie (0,683; 0,776; 0,763; 0,748)
	o 3 kwartały	Piekary Śląskie (0,659; 0,639; 0,617)		o 3 kwartały	–
	o 4 kwartały	–		o 4 kwartały	Jaworzno (0,531; 0,619), Katowice (0,512; 0,482)
	o 5 kwartałów	–		o 5 kwartałów	–

cd. tabeli 5

1	2	3	4	5	6
Bielsko-Biala	bez opóźnień	–		bez opóźnień	Jaworzno (0,693; 0,691; 0,675; 0,657; 0,638; 0,616)
	o 1 kwartał	Piekary Śląskie (0,581; 0,559; 0,538; 0,639; 0,617), Tamowskie Góry (0,568; 0,545; 0,541; 0,514; 0,484)		o 1 kwartał	–
	o 2 kwartały	–	Rybnik	o 2 kwartały	Częstochowa (0,776; 0,754; 0,857; 0,848), Jastrzębie-Zdrój (0,650; 0,631; 0,726; 0,709), Tychy (0,617; 0,596; 0,572; 0,546)
	o 3 kwartały	–		o 3 kwartały	Katowice (0,648; 0,628; 0,605), Bielsko-Biala (0,771; 0,757; 0,742)
	o 4 kwartały	–		o 4 kwartały	Racibórz (0,510; 0,479)
	o 5 kwartałów	–		o 5 kwartałów	–
Częstochowa	bez opóźnień	Jastrzębie-Zdrój (0,587; 0,566; 0,545; 0,519; 0,491; 0,471)		bez opóźnień	Zawiercie (0,503; 0,0,577; 0,577; 0,554; 0,528; 0,498)
	o 1 kwartał	Bielsko-Biala (0,681; 0,769; 0,757; 0,746; 0,730), Zawiercie (0,576; 0,556; 0,532; 0,505; 0,474)		o 1 kwartał	–
	o 2 kwartały	Tamowskie Góry (0,540; 0,515; 0,511; 0,481)	Tamowskie Góry	o 2 kwartały	Żory (0,667; 0,649; 0,629; 0,606), Będzin (0,567; 0,543; 0,634; 0,611)
	o 3 kwartały	–		o 3 kwartały	Racibórz (0,543; 0,516; 0,487), Rybnik (0,533; 0,623; 0,600)
	o 4 kwartały	Żory (0,618; 0,596)		o 4 kwartały	–
	o 5 kwartałów	–		o 5 kwartały	Jastrzębie-Zdrój (0,609)

cd. tabeli 5

1	2	3	4	5	6
Jastrzębie-Zdrój	bez opóźnień	Częstochowa (0,587; 0,566; 0,545; 0,519; 0,491; 0,460)	Tychy	bez opóźnień	–
	o 1 kwartał	Będzin (0,500; 0,578; 0,666; 0,646; 0,624), Zawiercie (0,577; 0,555; 0,530; 0,503; 0,472)		o 1 kwartał	Katowice (0,561; 0,538; 0,513; 0,484; 0,452), Będzin (0,550; 0,527; 0,612; 0,589; 0,563), Racibórz (0,548; 0,525; 0,500; 0,588; 0,562), Żory (0,530; 0,506; 0,479; 0,449; 0,414)
	o 2 kwartały	Racibórz (0,764; 0,751; 0,737; 0,845)		o 2 kwartały	–
	o 3 kwartały	–		o 3 kwartały	–
	o 4 kwartały	Piekary Śląskie (0,739; 0,723)		o 4 kwartały	–
Jaworzno	o 5 kwartały	↑		o 5 kwartałów	–
	bez opóźnień	Rybnik (0,693; 0,691; 0,675; 0,657; 0,638; 0,616)	Zawiercie	bez opóźnień	Tarnowskie Góry (0,503; 0,577; 0,577; 0,554; 0,528; 0,498)
	o 1 kwartał	–		o 1 kwartał	–
	o 2 kwartały	Częstochowa (0,665; 0,647; 0,743; 0,727)		o 2 kwartały	Żory (0,683; 0,667; 0,647; 0,625), Jaworzno (0,582; 0,662; 0,642; 0,620)
	o 3 kwartały	Bielsko-Biała (0,869; 0,862; 0,854)		o 3 kwartały	–
	o 4 kwartały	Katowice (0,640; 0,618), Piekary Śląskie (0,630; 0,607)		o 4 kwartały	–
Mikolów	o 5 kwartałów	–		o 5 kwartały	Tychy (0,740)
	bez opóźnień	–	Żory	bez opóźnień	–
	o 1 kwartał	Zawiercie (0,489; 0,566; 0,653; 0,749; 0,734)		o 1 kwartał	–
	o 2 kwartały	–		o 2 kwartały	Jastrzębie-Zdrój (0,544; 0,519; 0,608; 0,583), Jaworzno (0,572; 0,652; 0,631; 0,608)
	o 3 kwartały	Jaworzno (0,641; 0,635; 0,612), Rybnik (0,558; 0,533; 0,504)		o 3 kwartały	–
	o 4 kwartały	Bielsko-Biała (0,521; 0,492)		o 4 kwartały	Racibórz (0,622; 0,598)
	o 5 kwartałów	Piekary Śląskie (0,614)		o 5 kwartałów	–

W przypadku Będzina zidentyfikowano 5 rynków podobnych, w tym 2 wyprzedzające (Racibórz z wyprzedzeniem o 1 kwartał i Piekary Śląskie z wyprzedzeniem o 3 kwartały) i 3 naśladowujące (Jastrzębie-Zdrój i Tychy z opóźnieniem o 1 kwartał, Tarnowskie Góry z opóźnieniem o 2 kwartały). Bielsko-Biała ma 6 rynków podobnych, w tym 2 wyprzedzające z wyprzedzeniem o 1 kwartał (Piekary Śląskie, Tarnowskie Góry) oraz 4 naśladowujące (z opóźnieniem o 1 kwartał – Częstochowa, z opóźnieniem o 3 kwartały – Jaworzno i Rybnik, z opóźnieniem o 4 kwartały – Mikołów). Dla Częstochowy ustalono 8 rynków podobnych, w tym 1 zbieżny (Jastrzębie-Zdrój), 4 wyprzedzające (o 1 kwartał – Bielsko-Biała, Zawiercie; o 2 kwartały – Tarnowskie Góry, o 4 – kwartały Żory) oraz 3 naśladowujące (z opóźnieniem o 2 kwartały – Jaworzno, Rybnik; z opóźnieniem o 5 kwartałów – Piekary Śląskie). W przypadku Jastrzębia-Zdroju występuje 9 rynków podobnych, w tym 1 zbieżny (Częstochowa), 4 wyprzedzające (o 1 kwartał – Będzin, Zawiercie; o 2 kwartały – Racibórz i o 4 kwartały – Piekary Śląskie) oraz 3 naśladowujące (z opóźnieniem o 2 kwartały – Rybnik, Żory; z opóźnieniem o 5 kwartałów – Tarnowskie Góry). Jaworzno wykazuje podobieństwo do 10 rynków lokalnych, w tym 1 jest zbieżny (Rybnik), 4 wyprzedzające (o 2 kwartały – Częstochowa; o 3 kwartały – Bielsko-Biała; o 4 kwartały – Katowice i Piekary Śląskie) oraz 4 naśladowujące (z opóźnieniem o 2 kwartały – Zawiercie i Żory, z opóźnieniem o 3 kwartały – Mikołów, z opóźnieniem o 4 kwartały – Racibórz). W przypadku Mikołowa wyodrębniono tylko rynki wyprzedzające (5), w tym: wyprzedzający o 1 kwartał (Zawiercie), o 3 kwartały (Jaworzno, Rybnik), o 4 kwartały (Bielsko-Biała) oraz o 5 kwartałów (Piekary Śląskie). Dla Piekar Śląskich zaobserwowano istnienie 11 rynków podobnych, w tym 5 wyprzedzających (o 2 kwartały – Katowice, Tychy; o 3 kwartały – Rybnik, Żory; o 5 kwartałów – Częstochowa) oraz 6 naśladowających (z opóźnieniem o 1 kwartał – Bielsko-Biała; o 2 kwartały – Racibórz; o 3 kwartały – Będzin; o 4 kwartały – Jaworzno, Jastrzębie-Zdrój; o 5 kwartałów – Mikołów). Dla Raciborza zidentyfikowano 9 rynków podobnych, w tym 3 wyprzedzające (o 2 kwartały – Piekary Śląskie, o 4 kwartały – Jaworzno i Katowice) oraz 6 naśladowających (z opóźnieniem o 1 kwartał – Będzin, Tychy; o 2 kwartały – Jastrzębie-Zdrój; o 3 kwartały – Tarnowskie Góry; o 4 kwartały – Rybnik, Żory). W przypadku Rybnika stwierdzono 10 rynków podobnych, w tym 1 zbieżny (Jaworzno), 6 wyprzedzających (o 2 kwartały: Częstochowa, Jastrzębie-Zdrój, Tychy; o 3 kwartały – Katowice, Bielsko-Biała; o 4 kwartały – Racibórz) oraz 3 naśladowujące (z opóźnieniem o 3 kwartały – Mikołów, Piekary Śląskie, Tarnowskie Góry). Dla Tarnowskich Gór zidentyfikowano 8 rynków podobnych, w tym 1 zbieżny (Zawiercie), 5 wyprzedzających (z wyprzedzeniem o 2 kwartały – Żory, Będzin; o 3 kwartały – Racibórz, Rybnik; o 5 kwartałów – Jastrzębie-Zdrój) oraz 2 naśladowujące (z opóźnieniem o 1 kwartał – Bielsko-Biała; o 2 kwartały – Częstochowa). Tychy wykazują podobieństwo do 7 ryn-

ków, w tym 4 wyprzedzają je (o 1 kwartał – Katowice, Będzin, Racibórz, Żory), a 3 naśladują (z opóźnieniem o 2 kwartały – Piekary Śląskie, Rybnik; o 5 kwartałów – Zawiercie). W przypadku Zawiercia ustalono 7 rynków podobnych, w tym 1 zbieżny (Tarnowskie Góry), 3 wyprzedzające (o 2 kwartały – Żory, Jaworzno; o 5 kwartałów – Tychy) oraz 3 naśladujące (z opóźnieniem o 1 kwartał: Częstochowa, Jastrzębie-Zdrój, Mikołów). Dla Żor zidentyfikowano 8 rynków podobnych, w tym 3 wyprzedzające (o 2 kwartały – Jastrzębie-Zdrój, Jaworzno; o 4 kwartały – Racibórz) oraz 5 naśladujących (z opóźnieniem o 1 kwartał – Tychy, o 2 kwartały – Tarnowskie Góry, Zawiercie; o 3 kwartały – Piekary Śląskie, o 4 kwartały – Częstochowa).

Podsumowanie

Uzyskane wyniki mogą służyć podmiotom działającym na analizowanych lokalnych rynkach nieruchomości (inwestorom, deweloperom) w podejmowaniu różnego typu decyzji. Informacje o podobieństwie rynków lokalnych pod względem zmiany cen w czasie przy występowaniu rynków naśladujących (opóźnionych) stanowią podstawę prognozowania z wykorzystaniem metod analogowych. Z tego względu sprawdzenie jakości tak zbudowanych prognoz i tym samym możliwości prognozowania cen na lokalnych rynkach nieruchomości mieszkaniowych za pomocą metod analogowych będzie stanowiło przedmiot dalszych badań.

Celem innych badań będzie porównanie podobieństw występujących między rynkami zlokalizowanymi w bliskiej odległości geograficznej (np. w jednym województwie) z podobieństwami występującymi między rynkami oddalonymi od siebie geograficznie, ale stanowiącymi duże rynki nieruchomości (miasta wojewódzkie). Kolejnym zagadnieniem będzie zweryfikowanie czy w każdym województwie rynkiem wiodącym jest miasto wojewódzkie. Prowadzona będzie także dalsza analiza występowania zjawiska konwergencji między rynkami lokalnymi.

Literatura

1. Barro R., Sala-i-Martin X. *Convergence*. „Journal of Political Economy” 1992, No. 100(2).
2. Berbeka J., *Konwergencja gospodarcza a konwergencja społeczna krajów Unii Europejskiej (15) w latach 1985-2002. Nierówności społeczne a wzrost gospodarczy*, Uniwersytet Rzeszowski, zeszyt nr 8, Rzeszów 2006.
3. Cieślak M., Jasiński R., *Miara podobieństwa funkcji*, „Przegląd Statystyczny” 1979, nr 3-4.

4. Dittmann I., *Statystyczna analiza średnich cen transakcyjnych 1 m² powierzchni mieszkaniowej na wybranych rynkach lokalnych w Polsce* (w druku).
5. Dittmann I., *Lokalne rynki mieszkaniowe w Polsce – podobieństwo pod względem zmian cen transakcyjnych oraz dostępności*, Referat na XX Krajową Konferencję Naukową Towarzystwa Naukowego Nieruchomości, Toruń, 28-30 maja 2012.
6. Dorius S.F., *Global demographic convergence? A reconsideration of changing intercountry inequality in fertility*, „Population and Development Review” 2008, No. 34(3).
7. Nykiel L., *Mechanizmy dynamiki i zróżnicowania cen na rynku mieszkaniowym*, Studia i Materiały Towarzystwa Naukowego Nieruchomości, 2007, Vol. 15, nr 3-4.
8. Próchniak M., Witkowski B., *Modelowanie realnej konwergencji w skali między-narodowej*, „Gospodarka Narodowa” 2006, nr 10.
9. *Raport o średnich cenach transakcyjnych metra kwadratowego lokalu mieszkalnego w wybranych miastach województwa śląskiego 2006-II kw. 2011*, Centrum AMRON.
10. Sala-i-Martin X., *The Classical Approach to Convergence Analysis*, „The Economic Journal” 1996, No. 106.
11. Soja E., *Konwergencja płodności krajów europejskich*, Studia Ekonomiczne”, Uniwersytet Ekonomiczny, Katowice (w druku).
12. Trojanek R., *Determinanty wahań cen na rynku mieszkaniowym*, Studia i Materiały Towarzystwa Naukowego Nieruchomości, 2008, vol. 19, nr 3.
13. Wójcik P., *Dywergencja czy konwergencja: dynamika rozwoju polskich regionów*, „Studia Regionalne i Lokalne” 2008, nr 2 (32).
14. Żelazowski K., *Regionalne zróżnicowanie cen i ich determinant na rynku mieszkaniowym w Polsce*, Studia i Materiały Towarzystwa Naukowego Nieruchomości, 2011, vol. 19, nr 3.

THE SIMILARITY OF THE CHANGES IN AVERAGE TRANSACTION PRICES OF 1 M2 OF HOUSING IN SELECTED CITIES OF SILESIA

Summary

The paper provides a statistical analysis of behaviour of average transaction prices of 1m² of apartment space on local markets in the Silesian Voivodeship. The conducted research reveals that time and space analogies occur on local apartment markets in this voivodeship, i.e. the similarity of behaviour of transaction price changes with time. The shape similarity measure was applied for evaluating the degree of similarity and identifying the amount of lags. It enabled distinguishing leading, convergent and imitating markets among the analysed apartment markets. There was analysed also if the phenomenon of sigma convergence or sigma divergence occurs among local markets. Sigma convergence in the period of 1Q2006 – 1Q2009 and the stabilization in the period of 2Q2009 – 2Q2011 was stated.

Daniel Iskra

Uniwersytet Ekonomiczny w Katowicach

WARTOŚĆ ZAGROŻONA OPCJI EUROPEJSKICH SZACOWANA PRZEDZIAŁOWO. SYMULACJE

Wprowadzenie

Jednym z aspektów współczesnej ekonomii jest zarządzanie ryzykiem związanym z prowadzoną inwestycją, czyli identyfikacja, pomiar oraz sterowanie ryzykiem. Przez ostatnie dziesięciolecia wprowadzono (proces ten cały czas się toczy) wiele miar potrafiących skwantyfikować wybrany rodzaj ryzyka. W przypadku ryzyka rynkowego jedną z popularniejszych miar ostatnich lat jest wartość zagrożona (wartość narażona na ryzyko, *Value at Risk* – *VaR*) [1; 4], która zostanie użyta w niniejszym opracowaniu.

Koncepcję wartości narażonej na ryzyko można zdefiniować następująco [6]: jest to miara określająca wielkość straty, której osiągnięcie lub przekroczenie w ustalonej chwili czasu t jest równe zadanemu z góry prawdopodobieństwu α :

$$P(S_0 - S_t \geq VaR(\alpha, t)) = \alpha \quad (1)$$

gdzie:

S_0, S_t – odpowiednio początkowa (deterministyczna) i końcowa (losowa) wartość procesu cen instrumentu,

t – ustalony horyzont inwestycji,

α – ustalony poziom tolerancji (poziom istotności) dla szacowanej wartości *VaR*.

W celu wyznaczenia wartości zagrożonej danego instrumentu finansowego należy znać rozkład jego cen, co w praktyce najczęściej wiąże się z oszacowaniem parametrów wybranego rozkładu na podstawie próbki za pomocą odpowiednich estymatorów. Ze względu na fakt, iż estymatory (używane do estymacji parametrów rozkładu) są zmiennymi losowymi, ich oszacowana wartość zależy od próbki użytej w estymacji. Fakt ten można także wykorzystać do konstrukcji przedziałów ufności, które pokryją wartość zagrożoną z zadanym poziomem ufności [2; 3].

W celach poglądowych przeprowadzono symulacje, na podstawie których szacowano prognozę punktową wartości zagrożonej oraz wyznaczano przedział, który pokrywał wartość zagrożoną dla wybranej opcji kupna i sprzedaży na WIG20.

1. Model Blacka-Scholesa

Klasyczny model Blacka-Scholesa [5] jest jednym z najczęściej używanych modeli do wyceny opcji europejskich. Przy jego konstrukcji autorzy poczynili pewne założenia, m.in. takie, że dynamika cen instrumentu bazowego opisana jest geometrycznym ruchem Browna o stałych parametrach:

$$\frac{dS_t}{S_t} = \mu dt + \sigma dW_t \quad (2)$$

gdzie:

S_t – wartość instrumentu,

μ – współczynnik dryfu,

σ – współczynnik zmienności,

dW_t – przyrost procesu Wienera,

dt – przyrost czasu.

Rozwiązaniem powyższego równania różniczkowego stochastycznego jest funkcja:

$$S_t = S_0 e^{\left(\mu - \frac{1}{2}\sigma^2\right)t + \sigma W_t} \quad (3)$$

gdzie:

S_0 – aktualna wartość instrumentu.

W swoich rozważaniach przyjęli także brak wszelkich kosztów transakcyjnych oraz podatków oraz że krótkoterminowa wolna od ryzyka stopa procentowa r jest stała.

Na podstawie powyższych założeń Black i Scholes wyprowadzili wzory na ceny europejskich opcji kupna [5]:

$$c = S_0 N(d_1) - Xe^{-rT} N(d_2) \quad (4)$$

oraz opcji sprzedaży:

$$p = Xe^{-rT}N(-d_2) - S_0N(-d_1) \quad (5)$$

$$d_1 = \frac{\ln(\frac{S_0}{X}) + (r + \frac{\sigma^2}{2})T}{\sqrt{\sigma^2 T}} \quad (6)$$

$$d_2 = d_1 - \sqrt{\sigma^2 T} \quad (7)$$

gdzie:

$N(d)$ – dystrybuanta standaryzowanej zmiennej o rozkładzie normalnym,

S_0 – cena akcji w chwili $t = 0$,

T – czas pozostający do wygaśnięcia kontraktu,

X – cena wykonania opcji,

r – intensywność oprocentowania.

Symulacje

Zgodnie z definicją, aby wyznaczyć wartość zagrożoną opcji dla horyzontu t musimy znać jej wartość początkową, która zależy od oszacowanego parametru σ i wartości S_0, X, T, r (w pracy zakładano, że stopa r jest znana) oraz znać rozkład wartości opcji w chwili t . Rozkład ten jest powiązany z rozkładem zmiennej losowej, jaką jest cena instrumentu bazowego S_t . Do wyznaczenia prognozy punktowej wartości zagrożonej opcji za pomocą symulacji należy zatem wyestymować parametry $(\mu - \frac{1}{2}\sigma^2)$ i σ geometrycznego ruchu Browna.

Na ich podstawie wygenerować rozkład ceny instrumentu bazowego na chwilę t , a następnie wyznaczyć cenę początkową opcji w czasie $t=0$ i rozkład ceny opcji w czasie t .

W przypadku dynamiki cen instrumentu bazowego opisanej geometrycznym ruchem Browna, logarytmiczne stopy zwrotu o okresie Δt mają rozkład normalny:

$$N((\mu - \frac{1}{2}\sigma^2)\Delta t, \sigma\sqrt{\Delta t}) \quad (8)$$

Wykorzystując metodę momentów można wyznaczyć wzory na estymatory $(\mu - \frac{1}{2}\sigma^2)$ i σ^2 :

$$\begin{cases} \sigma^2 \Delta t = D^2(\ln(\frac{S_{t+\Delta t}}{S_t})) \\ (\mu - \frac{1}{2}\sigma^2)\Delta t = E(\ln(\frac{S_{t+\Delta t}}{S_t})) \end{cases} \quad (9)$$

W tym celu należy podstawić do powyższego układu równań wzory na estymatory wartości oczekiwanej $E(\ln(\frac{S_{t+\Delta t}}{S_t}))$ i wariancji $D^2(\ln(\frac{S_{t+\Delta t}}{S_t}))$ rozkładu normalnego logarytmicznych stóp zwrotu.

Należy podkreślić, że w rzeczywistości wartość oczekiwana czy wariancja rozkładu nie jest zmienną losową, podobnie jak sama wartość zagrożona wielkości te są deterministyczne. W praktyce jednak nie znamy ich prawdziwych wartości, a jedynie je oszacowujemy z pewnym błędem estymacji. Aby uwzględnić opisywaną niepewność będzie wyznaczany przedział ufności, który pokryje wartość zagrożoną z zadaniem poziomem ufności. Przedział ten, podobnie jak w przypadku prognozy punktowej również będzie wyznaczany symulacyjnie.

Dla rozkładu normalnego $N(E(X), D(X))$ nieobciążone estymatory wariancji oraz wartości oczekiwanej są niezależne i powiązane ze zmiennymi losowymi [7]:

$$H = \frac{(n-1)\hat{D}^2(X)}{D^2(X)} \quad (10)$$

$$T = \frac{(\hat{E}(X) - E(X))\sqrt{n-1}}{\hat{D}(X)} \quad (11)$$

gdzie:

H – zmienna o rozkładzie chi-kwadrat z $n-1$ stopniami swobody,

T – zmienna o rozkładzie t-Studenta z $n-1$ stopniami swobody,

n – liczebność próbek.

Jak już wcześniej zaznaczono, wielkości σ^2 i $(\mu - \frac{1}{2}\sigma^2)$ nie są zmiennymi losowymi, dlatego wprowadzono zmienne losowe $(\sigma^2)^*$, $(\mu - \frac{1}{2}\sigma^2)^*$, których rozkłady powinny pokryć odpowiednio opisywane wielkości.

Na podstawie wzorów (9), (10) i (11) można wyznaczyć rozkłady zmiennych losowych $(\sigma^2)^*$, $(\mu - \frac{1}{2}\sigma^2)^*$ (pod warunkiem, że znane są wielkości $\hat{E}(X)$, $\hat{D}^2(X)$):

$$(\sigma^2)^* = D^2(X) \frac{1}{dt} = \frac{(n-1)\hat{D}^2(X)}{H} \frac{1}{dt} \quad (12)$$

$$(\mu - \frac{1}{2}\sigma^2)^* = E(X) \frac{1}{dt} = (\hat{E}(X) - \frac{T\hat{D}(X)}{\sqrt{n-1}}) \frac{1}{dt} \quad (13)$$

gdzie:

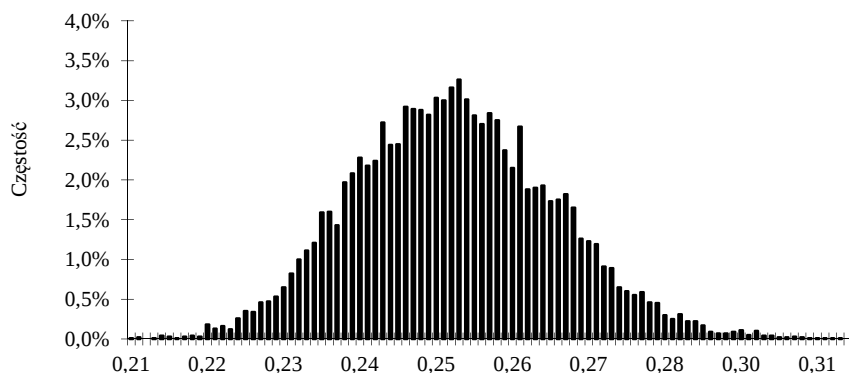
$\hat{E}(X)$, $\hat{D}^2(X)$ – wyestymowane wielkości wartości oczekiwanej i wariancji,
 H , T – zmienne losowe o rozkładzie chi-kwadrat i t -Studenta, oba z $n-1$ stopniami swobody.

Wygenerowane rozkłady zmiennych $(\sigma^2)^*$ oraz $(\mu - \frac{1}{2}\sigma^2)^*$ (na podstawie wzorów (12) i (13)) pokryją odpowiednio ich prawdziwe wartości.

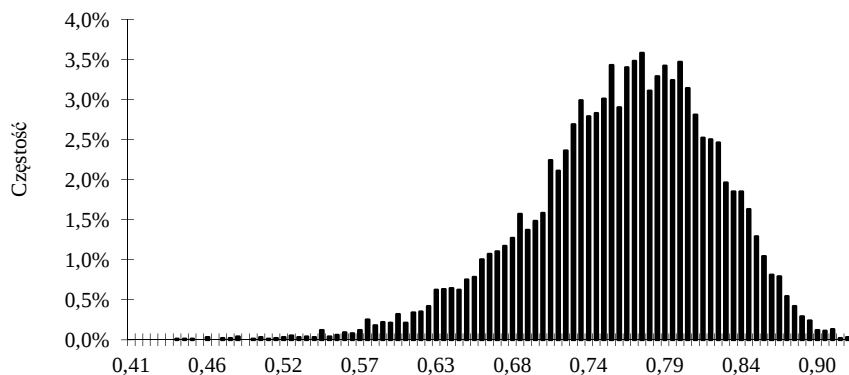
Dla prognozy punktowej rozkład zmiennej S_t generowano na podstawie 10 000 realizacji, estymując wcześniej parametry $(\mu - \frac{1}{2}\sigma^2)$ i σ^2 ze 100 jednodniowych logarytmicznych stóp zwrotu (por. wzór (9)). W przypadku prognozy przedziałowej parametry σ^2 i $(\mu - \frac{1}{2}\sigma^2)$ wyznaczano na podstawie 10 000 losowych realizacji zmiennej H oraz T (por. wzory (11) i (12)), a następnie dla każdej z 10 000 par tych parametrów generowano rozkład S_t również na podstawie 10 000 losowych realizacji (na podstawie wzoru (3)). Wartość oczekiwaną i wariancję logarytmicznych stóp zwrotu estymowano na podstawie próbki stuelementowej, natomiast stopę r obliczano jako średnią ważoną stóp WIBOR (wagi były ilorazem okresu stopy do czasu wygaśnięcia opcji).

Na rys. 1 i 2 przedstawiono histogramy zmiennej opisującej jednodniową oraz 15-dniową względną wartość zagrożoną (wartość zagrożoną w stosunku do początkowej wartości opcji: $\frac{VaR}{p_0}$, która będzie nazywana także $wzVaR$, do-

datnia wartość oznacza spadek wartości opcji) opcji OW20U1250. Jest to opcja sprzedaży na WIG20 o dacie zapadalności na wrzesień 2011 r. i wartości wykonania 2500. Poziom tolerancji VaR ustalono jako $\alpha = 0,05$. Stopa procentowa $r = 0,04$. Rozkłady generowano na dzień 30.03.2011 r., zatem przedziały odczytane z wygenerowanych rozkładów jako $[Q_{0,025}; Q_{0,975}]$, (Q_p – kwantyl rzędu p) powinny pokryć wartość zagrożoną odpowiednio na dzień następny i na 15 dni z poziomem ufności 0,95.



Rys. 1. Histogram względnej jednodniowej wartości zagrożonej

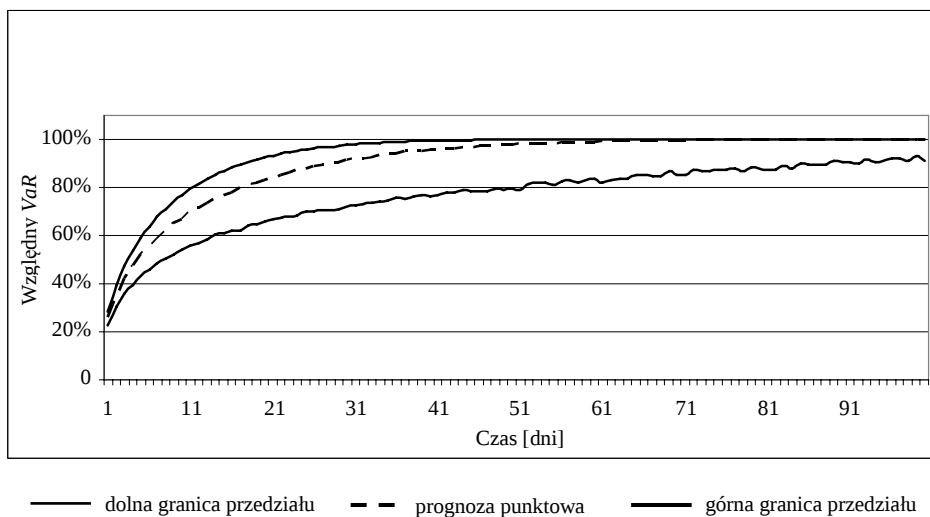


Rys. 2. Histogram względnej 15-dniowej wartości zagrożonej

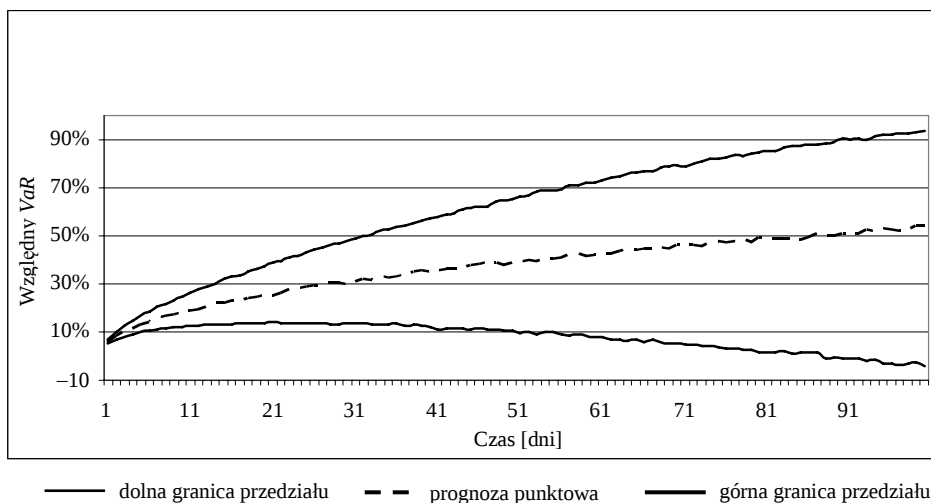
W przypadku jednodniowej wartości zagrożonej odczytany przedział $[0,23; 0,28]$ powinien pokryć wartość zagrożoną opcji z poziomem ufności 0,95. Jednodniowa prognoza punktowa wyniosła 0,26 (w tym przypadku parametry $(\mu - \frac{1}{2}\sigma^2)$ i σ^2 wynoszą odpowiednio (co do zaokrąglenia) 0,07 i 0,02 w skali roku). Rozpiętość przedziału ufności w stosunku do prognozy punktowej jest równa 22%.

Dla prognozy 15-dniowej odczytany przedział to $[0,61; 0,87]$, prognoza punktowa 0,78 (parametry nie ulegają zmianie), natomiast rozpiętość tego przedziału do prognozy punktowej jest równa 33%. Jak można zauważyć, w tym przypadku, stawiając prognozę punktową można się pomylić nawet o 15% jej wartości w wyniku błędu estymacji parametrów modelu.

Poniżej przedstawiono wykres wzVaR o coraz dłuższym horyzoncie VaR dla opcji sprzedaży OW20U1250 oraz opcji kupna OW20I1230 wyznaczanego w dacie 30.03.2011 r. (opcje na WIG20). Począwszy od wzVaR jednodniowego skończywszy na wzVaR 100-dniowym. Poziom tolerancji wartości zagrożonej był równy $\alpha = 0,05$. Przedziały pokrywały względną wartość zagrożoną z poziomem ufności 0,95.



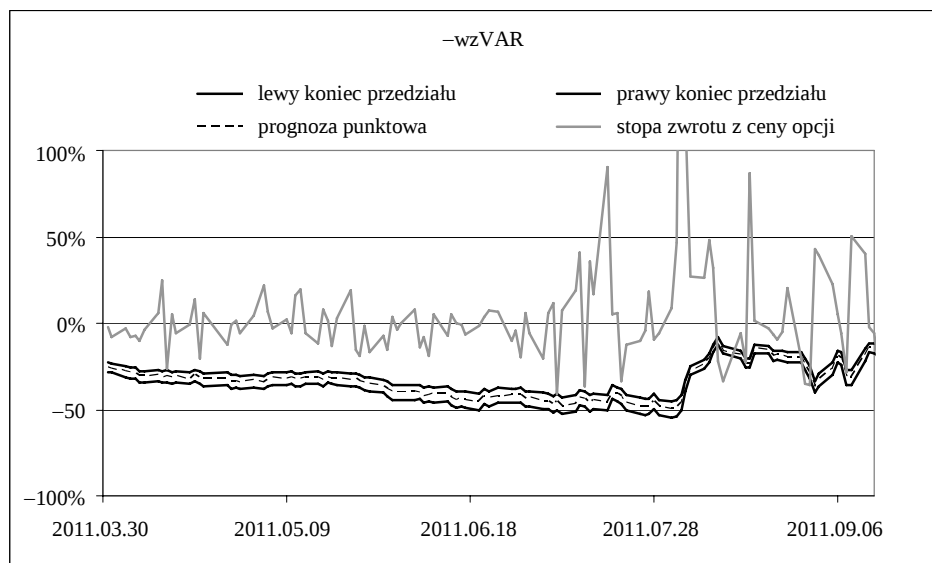
Rys. 3. Prognozy względnej wartości zagrożonej dla opcji OW20U1250 (opcja sprzedaży na WIG20 z terminem wykonania we wrześniu 2011 r. i wartością wykonania 2500)



Rys. 4. Prognozy względnej wartości zagrożonej dla opcji OW20I1230 (opcja kupna na WIG20 z terminem wykonania we wrześniu 2011 r. i wartością wykonania 2300)

Na wykresach można obserwować, jak zmienia się prognoza punktowa wartości zagrożonej wraz z wydłużaniem się czasu, dla którego była wyznaczana oraz jaki wpływ na dokładność szacowania ma estymacja parametrów modelu. Wraz z wydłużaniem się horyzontu wartości zagrożonej zwiększa się rozpiętość jej przedziału ufności aż górna granica osiągnie poziom 100%. Dla odpowiednio długiego horyzontu wartości zagrożonej dolna granica przedziału ufności też będzie dążyć do 100% w przypadku opcji sprzedaży i do 0 w przypadku opcji kupna. Taka sytuacja wynika z faktu, że wraz z wydłużaniem czasu do wygaśnięcia opcji prawdopodobieństwo, że wartość indeksu w dacie wykonania będzie w przedziale $[0; X]$ zmierza do zera (dla parametrów estymowanych na dzień 30.03.2011 r.).

Na kolejnym rysunku przedstawiono wyniki symulacji, w wyniku której szacowano jednodniowe względne wartości zagrożone z poziomem tolerancji 0,05 zarówno dla prognozy przedziałowej, jak i punktowej. Symulacje przeprowadzono dla opcji OW20U1250 zmniejszając za każdym razem czas pozostały do jej wygaśnięcia o jeden dzień. Względne wartości zagrożone porównano z realnymi cenami tej opcji, a dokładniej z ich stopami zwrotu. Pamiętając, że dodatnia wartość zagrożona oznacza spadek ceny opcji, na wykresie przedstawiono względne wartości zagrożone przemnożone przez (-1) (oznaczymy tę wartość jako $wzVaR$).



Rys. 5. Jednodniowe prognozy wartości zagrożonej

Minimalna zaobserwowana rozpiętość przedziału ufności w stosunku do prognozy punktowej $wzVaR$ dla tego samego dnia wyniosła około 18%, maksymalna około 40%. Średnia rozpiętość przedziału ufności w stosunku do prognozy punktowej jest równa 23,3%, natomiast odchylenie standardowe około 5%. Widać zatem, że prognoza punktowa względnej wartości zagrożonej może w niektórych przypadkach znacznie różnić się od wartości realnej.

Podsumowanie

W przypadku prognoz punktowych wartości zagrożonej szansa na jej bezbłędne oszacowanie jest znikoma (dokładniej równa 0). Oczywiście błąd estymacji może być tak mały, że różnice pomiędzy wyznaczoną a rzeczywistą wartością zagrożoną będą nieistotne i taka prognoza będzie skuteczna. Istnieje jednak alternatywa w postaci estymacji przedziałowej, w której można wyznaczyć przedział pokrywający wartość zagrożoną z góry zadany poziom ufności. W opracowaniu przedstawiono próbę wyznaczania wartości zagrożonej wybranych opcji europejskich na WIG20 za pomocą symulacji. Na ich podstawie można zaobserwować, jak zmienia się rozpiętość prognoz przedziałowych w czasie, czy też w zależności od horyzontu wartości zagrożonej. W bardzo prosty sposób obrazują one, jakie błędy mogą wystąpić podczas estymacji punktowej wartości zagrożonej.

Literatura

1. Alexander C., *Market Risk Analysis: Value at Risk Models*, Vol. IV, John Wiley & Sons, England 2008.
2. Contreras P., Satchell S., *A Bayesian Confidence Interval for Value-at-Risk* 2003, <http://www.dspace.cam.ac.uk/handle/1810/380>
3. Dowd K., *Assessing VaR Accuracy*, 2000 <http://www.smartquant.com/references/VaR/var14.pdf>
4. Holton G.A., *Value at Risk. Theory and Practice*, Academic Press, USA 2003,
5. Hull J., *Kontrakty terminowe i opcje. Wprowadzenie*, WIG-PRESS, Warszawa 1999.
6. Jorion P., *Value at risk: the new benchmark for managing financial risk*, McGraw-Hill 2001.
7. Sobczyk M., *Statystyka. Aspekty praktyczne i teoretyczne*, UMCS, Lublin 2006.

VaR ESTIMATED INTERVAL EUROPEAN OPTIONS. SIMULATIONS

Summary

This paper presents a procedure for determining the value at risk ranges covering European options with a given level of confidence. Interval forecast VaR takes into account the uncertainty associated with the estimation error of the model parameters used. Option pricing model adapted Black-Sholes, and studies based on simulations.

Jan Acedański

Uniwersytet Ekonomiczny w Katowicach

KRYTERIA WYBORU DYNAMICZNYCH MODELI CZYNNIKOWYCH DLA CELÓW PROGNOSTYCZNYCH

Wprowadzenie

Dynamiczne modele czynnikowe (*dynamic factor models* – DFM) są ważnym narzędziem wykorzystywanym w krótkookresowym prognozowaniu makroekonomicznym. Podejście to często jest stosowane przez banki centralne jako narzędzie wspomagające krótkookresowe prognozowanie kluczowych zmiennych makroekonomicznych [5; 6; 7]. Istota metody sprowadza się do agregacji dużej liczby potencjalnych zmiennych objaśniających do kilku wzajemnie niezależnych czynników, które następnie wykorzystywane są do prognozowania wybranej zmiennej. Wyodrębnianie czynników najczęściej dokonywane jest za pomocą klasycznej metody składowych głównych lub jej modyfikacji. Równanie progностyczne opisujące zależność pomiędzy zmienną prognozowaną a czynnikami ma zwykle postać liniową. Oprócz czynników w równaniu tym mogą wystąpić ich opóźnienia, a także składniki o charakterze autoregresyjnym.

Jednym z ważniejszych problemów związanych ze stosowaniem omawianego podejścia jest specyfikacja modelu. W szczególności konieczne jest podjęcie decyzji dotyczących liczby czynników branych pod uwagę przy prognozowaniu, a także liczby ewentualnych opóźnień występujących w równaniu progностycznym. W literaturze przedmiotu opisywanych jest wiele podejść stosowanych w celu wyboru modelu dla celów progностycznych, np. wybór optymalnej liczby czynników dokonywany był często na podstawie zmodyfikowanych kryteriów informacyjnych zaproponowanych przez [3; 7], a wybór dokładnej specyfikacji równania progностycznego był oparty na standardowym kryterium informacyjnym BIC [2; 13]. Dość często także liczba czynników oraz opóźnień ustalana była *ad hoc* na arbitralnym, zwykle niewielkim poziomie [5; 12].

W ostatnich latach pojawiły się w literaturze podejścia, które pozwalają na dokonanie jednoczesnego wyboru liczby czynników oraz opóźnień w równaniu prognostycznym. Bai i Ng [3] oraz Groen i Kapetanios [11] zaproponowali zastosowanie w tym celu zmodyfikowanych kryteriów informacyjnych. Jeszcze inną możliwością jest wybór najlepszego modelu na podstawie analizy jakości prognoz wygasłych modeli z różną liczbą czynników oraz opóźnień.

Celem pracy jest porównanie omówionych trzech podejść do wyboru najlepszego modelu DFM dla celów prognostycznych: zmodyfikowanych kryteriów informacyjnych, wyłącznej analizy dokładności prognoz wygasłych oraz podejścia tradycyjnego łączącego kryteria informacyjne i analizę prognoz wygasłych. Badane jest, które z tych trzech podejść generuje najlepsze prognozy poza próbą. Podstawowym kryterium oceny jakości prognoz jest błąd średniokwadratowy RMSE. W badaniach wykorzystano zarówno szeregi generowane na podstawie symulacji Monte Carlo, jak i dane rzeczywiste dotyczące prognozowania miesięcznej stopy inflacji w Polsce.

1. Modele czynnikowe w prognozowaniu

Istota prognozowania za pomocą dynamicznych modeli czynnikowych opiera się na założeniu, że źródłem dynamiki bardzo wielu zmiennych ekonomicznych jest niewielka liczba niezależnych czynników [16]. Jeżeli $\mathbf{X}_t$ oznacza wektor N zmiennych objaśniających obserwowanych w chwili t , wtedy powyższe założenie można sformalizować w postaci:

$$\mathbf{X}'_t = \mathbf{F}'_t \mathbf{L} + \boldsymbol{\varepsilon}'_t, \quad t = 1, 2, \dots, T \quad (1)$$

gdzie $\mathbf{F}_t$ jest K -wymiarowym wektorem czynników, $\mathbf{L}$ – $K \times N$ -wymiarową macierzą ładunków czynnikowych, natomiast $\boldsymbol{\varepsilon}_t$ reprezentuje N -wymiarowy wektor składników losowych.

O czynnikach zakłada się, że są one między sobą niezależne: $E(F_{it}, F_{jt}) = 0$ dla $i \neq j$ oraz niezależne od zaburzeń losowych: $E(F_{it}, \varepsilon_{it}) = 0$. Ponadto zakłada się, że składniki losowe mają zerową wartość oczekiwaną $E(\varepsilon_{it}) = 0$ oraz są niezależne między sobą $E(\varepsilon_{it}, \varepsilon_{jt}) = 0$, $i \neq j$. Oprócz powyższych założeń zwykle przyjmuje się także, że macierz ładunków czynnikowych jest stała oraz że nie występuje autokorelacja składników losowych, chociaż uchylenie tych założeń jest możliwe [9]. Liczba czynników K oraz ich realizacje nie są obserwowane i muszą być estymowane przez badacza. Fakt, że czynniki nie są obserwowane sprawia, że w omawianej sytuacji nie powinno się stosować standardowych kryteriów informacyjnych przy wyborze liczby czynników oraz opóźnień w równaniu prognostycznym.

Najpopularniejszą metodą estymacji czynników jest metoda składowych głównych (*principal components analysis* – PCA). Pozwala ona na znalezienie takich N -wymiarowych wektorów czynników $\hat{\mathbf{F}}_t$ oraz takiej $N \times N$ -wymiarowej macierzy ładunków czynnikowych $\hat{\mathbf{L}}$, że:

$$\mathbf{X}'_t = \hat{\mathbf{F}}'_t \hat{\mathbf{L}}, \quad t = 1, 2, \dots, T \quad (2)$$

Estymowane czynniki $\hat{\mathbf{F}}_t$ są wzajemnie niezależne oraz uporządkowane w takiej kolejności, że zmiany kolejnego czynnika wyjaśniają coraz mniejszy odsetek wahań zmiennych obserwowanych z wektorów $\mathbf{X}_t$.

Przy szacowaniu $\hat{\mathbf{F}}_t$ oraz $\hat{\mathbf{L}}$ metodą składowych głównych wykorzystuje się dekompozycję macierzy $\mathbf{X}\mathbf{X}'$, w której $\mathbf{X} = [\mathbf{X}_1 \ \mathbf{X}_2 \ \dots \ \mathbf{X}_T]'$. Kolejne kolumny macierzy $\hat{\mathbf{L}}$ są wektorami własnymi macierzy $\mathbf{X}\mathbf{X}'$ związanymi z kolejnymi, uporządkowanymi malejąco wartościami własnymi tej macierzy [16]. Przy danej macierzy ładunków czynnikowych realizacje czynników można łatwo obliczyć jako $\hat{\mathbf{F}}'_t = \mathbf{X}'_t \hat{\mathbf{L}}^{-1}$. Czasami zdarza się, że niektóre wektory $\mathbf{X}_t$ są niekompletne, szczególnie na końcu lub początku próby. W takiej sytuacji przy estymacji czynników stosowany jest dodatkowo algorytm maksymalizacji oczekiwań (*expectations maximization* – EM) [16]. Algorytm ten ma charakter rekurencyjny, składający się z dwóch podstawowych kroków. W pierwszym kroku w oparciu o dane wartości $\hat{\mathbf{F}}_t$ oraz $\hat{\mathbf{L}}$ uzupełniane są brakujące elementy w niekompletnych wektorach $\mathbf{X}_t$. W drugim kroku korzystając już z kompletnych wektorów $\mathbf{X}_t$ dokonuje się estymacji $\hat{\mathbf{F}}_t$ oraz $\hat{\mathbf{L}}$ zwykłą metodą składowych głównych. Startowe wartości $\hat{\mathbf{F}}_t$ oraz $\hat{\mathbf{L}}$ w procedurze wyznacza się najczęściej biorąc pod uwagę tylko kompletne wektory $\mathbf{X}_t$.

Metoda PCA zakłada statyczną zależność pomiędzy zmiennymi objaśnianymi a czynnikami. W przypadku przyjęcia założenia, że zależność ta nie ma charakteru statycznego stosowane są inne metody estymacji czynników, np. wykorzystujące analizę spektralną [9] lub filtry Kalmana [8].

W drugim kroku, po wyborze właściwej liczby czynników specyfikowane jest równanie prognostyczne, które można przedstawić w następującej postaci:

$$Y_{t+h} = \alpha_0 + \sum_{i=0}^{ny} \alpha_{i+1} Y_{t-i} + \sum_{j=0}^{nf} \beta_j \hat{\mathbf{F}}_t^{(k)} + \varepsilon_{Yt} \quad (3)$$

gdzie Y_t oznacza zmienną prognozowaną, $\hat{\mathbf{F}}_t^{(k)}$ reprezentuje wektor k pierwszych oszacowanych czynników, ε_{Yt} jest składnikiem losowym, natomiast α_i oraz β_j są parametrami równania. Horyzont prognozy oznaczony został przez h , n_y jest nieznaną liczbą opóźnionych wartości zmiennej prognozowanej, natomiast n_f reprezentuje opóźnienia czynników. Parametry równania (3) szacowane są metodą najmniejszych kwadratów.

Należy zwrócić uwagę, że omówione tutaj podejście jest zbliżone do koncepcji prognozowania za pomocą wskaźników wyprzedzających. Równanie (3) wskazuje bowiem, że czynniki mają charakter wskaźników wyprzedzających. W tym miejscu ujawnia się pewna niespójność całego omawianego podejścia, gdyż przy estymacji czynników zakładano, że zależności pomiędzy czynnikami a zmiennymi $\mathbf{X}_t$ mają charakter jednoczesny. Tymczasem w praktyce do zbioru zmiennych objaśniających zalicza się bardzo wiele zmiennych nie patrząc na to, czy mają one charakter wskaźników wyprzedzających względem zmiennej prognozowanej. Specyfikacja (3) ma więc charakter *ad hoc*.

Alternatywne podejście polega na specyfikacji równania opisującego zależność jednoczesne pomiędzy Y_t a $\hat{\mathbf{F}}_t^{(k)}$ wraz z ewentualnymi opóźnieniami [7], jednak w takiej sytuacji stawianie prognoz wymaga znajomości przyszłych wartości czynników, a więc konieczne jest skonstruowanie jeszcze jednego modelu służącego do prognozowania przyszłych wartości czynników. W tym celu najczęściej stosowane są modele wektorowej autoregresji.

2. Metody wyboru liczby czynników oraz opóźnień w równaniu prognostycznym

W literaturze opisywanych jest kilka metod wyboru optymalnej liczby czynników oraz specyfikacji równania prognostycznego. Można je podzielić na dwie grupy. W pierwszej obie decyzje są dokonywane oddzielnie. Do drugiej grupy należą metody jednoczesnego wyboru liczby czynników i opóźnień.

2.1. Kryteria Bai i Ng wyboru liczby czynników

Najpopularniejszymi kryteriami wyboru liczby czynników są modyfikacje standardowych kryteriów informacyjnych zaproponowane przez Bai i Ng [3].

Metoda składowych głównych generuje N czynników, a więc tyle samo, ile jest zmiennych obserwowanych tworzących wektory $\mathbf{X}_t$. Bai oraz Ng zaproponowali kryteria, które pozwalają na estymację nieznaną liczbę K czynników

rzeczywistych. Metoda ta daje właściwe oszacowania asymptotycznie, tzn., gdy badacz dysponuje jednocześnie długimi szeregami czasowymi – $T \rightarrow \infty$ oraz dużym zbiorem zmiennych objaśniających – $N \rightarrow \infty$. W omawianym podejściu rozważa się sekwencję modeli postaci (1), w których uwzględnianych jest k pierwszych czynników $k = 1, 2, \dots, k_{MAX}$ wraz z odpowiednimi postaciami macierzy ładunków czynnikowych uzyskanych metodą głównych składowych. Jeżeli przez $V(k)$ oznaczyć łączną wariancję wszystkich reszt w modelu z k czynnikami:

$$V(k) = \frac{1}{NT} \sum_{i=1}^N \sum_{t=1}^T \left(X_{it} - \sum_{j=1}^k \hat{F}_{jt} \hat{L}_{ji} \right)^2 \quad (4)$$

natomiast $\hat{\sigma}^2$ oznacza oszacowanie wariancji wszystkich składników losowych ε_i , wtedy optymalna liczba czynników $\hat{k}$ powinna zostać ustalona na takim poziomie, aby wartości kryteriów:

$$PC(k) = V(k) + k \hat{\sigma}^2 \frac{N+T}{NT} \ln \left(\frac{NT}{N+T} \right) \quad (5)$$

$$IC(k) = \ln(V(k)) + k \hat{\sigma}^2 \frac{N+T}{NT} \ln(\min\{N, T\}) \quad (6)$$

były jak najmniejsze. W praktycznych zastosowaniach $\hat{\sigma}^2$ może być przybliżane przez $V(k_{MAX})$. W pracy Bai i Ng powyższe kryteria były oznaczane jako $PC_{p1}(k)$ oraz $IC_{p2}(k)$. W badaniach porównawczych kryteria te odznaczały się dobrymi własnościami w porównaniu do innych kryteriów rozważanych w cytowanym opracowaniu. Należy jeszcze zaznaczyć, że w ostatnich latach pojawiło się w literaturze kilka alternatywnych kryteriów o korzystnych własnościach w porównaniu do kryteriów przedstawionych powyżej, choć nie są one jeszcze tak popularne w zastosowaniach prognostycznych [1; 14; 15].

2.2. Metody analizy prognoz wygasłych w wyborze liczby opóźnień w równaniu prognostycznym

Wyboru właściwej liczby opóźnień n_y i n_f można także dokonywać poprzez porównanie jakości prognoz wygasłych modeli z różnymi liczbami opóźnień. W zbiorze T obserwacji wyodrębnia się r okien obserwacyjnych, poczynwszy od pierwszego obejmującego obserwacje o numerach 1, 2, ..., T_p . Kolejne okna są

albo przesuwane o jedną obserwację wprzód (*rolling window*), albo rozszerzane o kolejną obserwację (*expanding window*). W pierwszym przypadku okna liczą zawsze obserwacji T_p . W drugim liczba obserwacji wzrasta o jeden.

Następnie dla każdej metody m oraz okna i szacowane są parametry równania (3) oraz wyznaczane są prognozy na h okresów do przodu: $\hat{Y}_{T_p}(m, h, i)$, które są porównywane z rzeczywistymi wartościami zmiennej prognozowanej w danym okresie. Jako podstawową miarę jakości prognoz przyjęto w pracy błąd średniokwadratowy:

$$RMSE(m, h) = \sqrt{\frac{1}{r} \sum_{i=1}^r \left(\hat{Y}_{T_p}(m, h, i) - Y_{T_p+i+h} \right)^2} \quad (7)$$

Jako najlepszy model do prognozowania przy danym horyzoncie prognozy wybierano ten model, który cechował się najmniejszym błędem średniokwadratowym.

W literaturze przedmiotu cały czas toczy się dyskusja dotycząca właściwego sposobu konstrukcji okien prognostycznych. Za oknami rozszerzanymi przemawia fakt, że pozwalają one uwzględnić wszystkie obserwacje. Zwolennicy okien przesuwanych podkreślają natomiast, że podejście to cechuje się większą odpornością w przypadku niestabilności procesu generującego dane. Badania empiryczne i symulacyjne nie są w stanie jednoznacznie rozstrzygnąć tej kwestii [10].

Biorąc pod uwagę dwa kryteria wyboru optymalnej liczby czynników oraz dwie metody konstrukcji okien przy wyborze najlepszej specyfikacji równania prognostycznego łącznie w pracy rozważane były cztery tradycyjne podejścia do wyboru najlepszego modelu DFM dla celów prognozowania.

2.3. Metody jednoczesnego wyboru liczby czynników i opóźnień

Na alternatywne podejścia do wyboru optymalnego modelu DFM, mające na celu jednoczesne podjęcie decyzji dotyczącej liczby czynników oraz opóźnień w równaniu prognostycznym można patrzeć jako na rozszerzenia metod opisanych w poprzednich dwóch rozdziałach. Można je podzielić na dwie grupy. W pierwszej znajdują się metody wykorzystujące zmodyfikowane wersje kryteriów informacyjnych. Druga bazuje na analizie prognoz wygasłych.

Odnosnie do kryteriów z pierwszej grupy Bai i Ng [4] zaproponowali kryterium finalnego błędu predykcji (*final prediction error* – FPE), które przy dużych próbach pozwala na wybór modelu postaci (3) minimalizującego średniokwadratowy błąd predykcji przy danym horyzoncie prognozy:

$$FPE = \ln(\hat{\sigma}^2) + (2 + n_y + k(1 + n_f)) \frac{\ln T}{T} + \frac{\hat{\gamma}' \hat{\Sigma} \hat{\gamma}}{\hat{\sigma}^2} \frac{\ln N}{N} \quad (8)$$

gdzie $\hat{\gamma}$ oznacza wektor ocen wszystkich parametrów strukturalnych α_i , β_j w równaniu (3), natomiast $\hat{\Sigma}$ reprezentuje oszacowanie macierzy kowariancji zmiennych wszystkich zmiennych niezależnych równania (3).

Do pierwszej grupy należą także kryteria podane przez Groena i Kapetaniosa [11]. Są one modyfikacją standardowych kryteriów informacyjnych uwzględniającą fakt, że zmienne objaśniające w równaniu prognostycznym nie są z góry ustalone, ale muszą być estymowane. W odniesieniu do równania (3) kryteria te przyjmują następujące formy:

$$BICM = \frac{T}{2} \ln(\hat{\sigma}^2) + (2 + n_y) \ln T + k(1 + n_f) \ln T \left(1 + \frac{T}{N}\right) \quad (9)$$

$$HQICM = \frac{T}{2} \ln(\hat{\sigma}^2) + 2(2 + n_y) \ln(\ln T) + 2k(1 + n_f) \ln(\ln T) \left(1 + \frac{T}{N}\right) \quad (10)$$

We wszystkich trzech przypadkach jako najlepszy model wybiera się ten, dla którego wartości kryteriów są najmniejsze.

Wyboru najlepszego modelu prognostycznego w całości można również dokonać porównując jakość prognoz wygasłych, zgodnie z procedurą opisaną w poprzednim rozdziale, przy czym oprócz liczby opóźnień w równaniu (3) uwzględnia się teraz dodatkowo różne liczby czynników w modelu n_f .

3. Metodologia badań porównawczych omawianych metod

Łącznie w pracy porównywano 9 metod wyboru najlepszego modelu czynnikowego dla celów prognostycznych. Były to 4 podejścia klasyczne łączące kryteria Bai-Ng z badaniem jakości prognoz wygasłych – w dalszej części pracy oznaczane jako: *PC_rol* (kryterium *PC* i okna przesuwane), *PC_exp* (kryterium *PC* i okno rozszerzane), *IC_rol* (kryterium *IC* i okna przesuwane), *IC_exp* (kryterium *IC* i okno rozszerzane) – oraz 5 podejść nowych pozwalających na jednoczesny wybór liczby czynników i liczby opóźnień w równaniu prognostycznym: zmodyfikowane kryteria informacyjne – *FPE*, *BICM*, *HQICM* oraz kryteria bazujące na badaniu jakości prognoz wygasłych – *RF_rol* (z oknami przesuwanymi) i *RF_exp* (z oknami rozszerzanymi).

Badania przeprowadzono zarówno na symulowanych, jak i na rzeczywistych szeregach czasowych. W przypadku danych symulowanych w pierwszym kroku generowano szereg czasowy K czynników:

$$\mathbf{F}_t = \rho_F \mathbf{F}_{t-1} + \mathbf{v}_{Ft}, \quad \mathbf{v}_{Ft} \sim N(\mathbf{0}, \mathbf{I}_K) \quad (11)$$

gdzie ρ_F określa jednakowy dla wszystkich czynników poziom autokorelacji, $\mathbf{I}_K$ symbolizuje natomiast K -wymiarową macierz jednostkową. Dla każdej symulacji liczbę czynników K wybierano losowo ze zbioru $\{1, 2, \dots, 5\}$. Następnie generowano szereg $(N+1) \times K$ -wymiarowych macierzy ładunków czynnikowych:

$$\begin{aligned} \mathbf{L}_t &= [1 + \psi_L (\boldsymbol{\eta}_{Lt} \otimes \mathbf{I}_K)] \mathbf{L}_{t-1} \\ \boldsymbol{\eta}_{Lt} &= \rho_L \boldsymbol{\eta}_{Lt-1} + \mathbf{v}_{Lt}, \quad \text{vec}(\mathbf{v}_{Lt}) \sim N(\mathbf{0}, \mathbf{I}_{(N+1)K}) \end{aligned} \quad (12)$$

przy czym ψ_L jest parametrem skalującym, $\boldsymbol{\eta}_{Lt}$ jest ciągiem $(N+1) \times K$ -wymiarowych realizacji procesu autoregresyjnego o współczynniku autokorelacji ρ_L , natomiast $\otimes$ reprezentuje iloczyn Kroneckera, a więc oznacza zamianę macierzy w wektor kolumnowy. Z powyższego zapisu wynika, że w ogólnym przypadku dopuszczano zmienność macierzy ładunków czynnikowych. W większości symulacji zakładano jednak, że $\psi_L = 0$, co oznaczało stałość macierzy $\mathbf{L}_t$. Początkowe wartości składowych wektora czynników $\mathbf{F}_0$ oraz macierzy ładunków czynnikowych $\mathbf{L}_0$ losowane były z rozkładu $N(0, 1)$. Wreszcie na podstawie wektorów $\mathbf{F}_t$ oraz macierzy $\mathbf{L}_t$ tworzone były $(N+1)$ -wymiarowe wektory obserwacji:

$$\mathbf{Z}_t = \mathbf{F}_t \mathbf{L}_t + \boldsymbol{\eta}_{Zt}, \quad \boldsymbol{\eta}_{Zt} = \rho_Z \boldsymbol{\eta}_{Zt-1} + \mathbf{v}_{Zt}, \quad \mathbf{v}_{Zt} \sim N(\mathbf{0}, \mathbf{I}_{N+1}) \quad (13)$$

Pierwsza kolumna macierzy $\mathbf{Z}_t$ reprezentowała zmienną prognozowaną Y_t . Pozostałe N kolumn traktowano jako potencjalne zmienne objaśniające $\mathbf{X}_t$. Przed prowadzeniem dalszych obliczeń kolumny macierzy $\mathbf{Z}_t$ zostały zstandaryzowane. Przyjmując powyższą specyfikację nie zakładano więc żadnego sprecyzowanego modelu parametrycznego dla zmiennej prognozowanej, ale traktowano ją jak jedną z wielu zmiennych ekonomicznych, których dynamika była kształtowana przez nieobserwowane czynniki.

Przeprowadzając badania przyjmowano, że prognozowanie nie ma charakteru jednorazowego, ale sekwencyjny. To znaczy zakładano, że prognozy stawiane są przez p kolejnych okresów zawsze na podstawie T ostatnich obserwacji. Takie postępowanie jest charakterystyczne dla omawianych modeli, które zwykle wykorzystywane są przez instytucje do regularnego comiesięcznego stawiania prognoz wybranych zmiennych. Symulowane szeregi czasowe liczyły więc $T + p + h_{MAX}$ obserwacji, gdzie h_{MAX} oznacza maksymalny horyzont prognozy. Na ich podstawie dla każdej metody m oraz dla każdego horyzontu prognozy h wybierano najlepszy model i generowano prognozy $\hat{Y}_{T+i-1}(m, h)$, $i = 1, 2, \dots, p$, które porównywano z rzeczywistymi realizacjami prognozowanej zmiennej. W ten sposób uzyskiwano ciąg p błędów prognoz:

$$u_i(m, h) = \hat{Y}_{T+i-1}(m, h) - Y_{T+i-1+h}, \quad i = 1, 2, \dots, p \quad (14)$$

Przy porównywaniu jakości prognoz uzyskanych różnymi metodami stosowano kryterium błędu średniokwadratowego prognoz wygasłych, przy czym aby ocenić, czy zaobserwowana w próbie p błędów różnica jakości była statystycznie istotna korzystano z testu warunkowej zdolności progностycznej (*conditional prediction ability* – CPA) opracowanego przez Giacomini i White'a [10]. Hipoteza podstawowa w omawianym teście wskazuje, że jakość prognoz generowanych przez dwa konkurencyjne podejścia była w próbie taka sama. Hipoteza alternatywna jest prostym zaprzeczeniem hipotezy badanej.

Jeżeli przez $\Delta L_i(m_1, m_2, h)$ oznaczyć różnicę między błędami średniokwadratowymi prognoz wygasłych o horyzoncie h uzyskanych z dwóch różnych modeli m_1 i m_2 dla próby i , natomiast przez S_i – s -elementowy wektor informacji dodatkowej dostępnej w próbie i , wtedy statystyka testowa testu CPA przyjmuje postać:

$$\chi^2 = p \bar{L}' \hat{\Omega}^{-1} \bar{L} \quad (15)$$

gdzie $\bar{L} = \frac{1}{p} \sum_{i=1}^p \Delta L_i(m_1, m_2, h) S_i$, natomiast $\hat{\Omega}$ oznacza oszacowanie macierzy kowariancji wektorów postaci $\Delta L_i(m_1, m_2, h) S_i$ uzyskane metodą Neweya-West.

Asymptotycznie, przy dużej liczbie porównywanych prognoz ($p \rightarrow \infty$), statystyka testowa ma rozkład χ^2 z s stopniami swobody. W przypadku braku warunkowej informacji dodatkowej $S_i = 1$. Wtedy test ma charakter bezwarunkowy. Test bezwarunkowy wskazuje podejście, które generowało średnio dokładniejsze prognozy. Można więc uznać, że powinno ono także wskazywać, które podejście będzie lepsze w pewnym, bliżej nieokreślonym momencie w przyszłości. Test względny odpowiada natomiast na pytanie, czy na podstawie innych dostępnych w próbie informacji, poza średnią jakością prognoz wygasłych, można wskazać podejście, które będzie generowało prognozy lepsze w określonych warunkach. Omawiany test może być stosowany tylko w sytuacji, gdy kolejne prognozy obliczane są na podstawie takiej samej liczby obserwacji, co jest dość istotnym ograniczeniem. Z tej właśnie przyczyny długość prób T była zawsze stała.

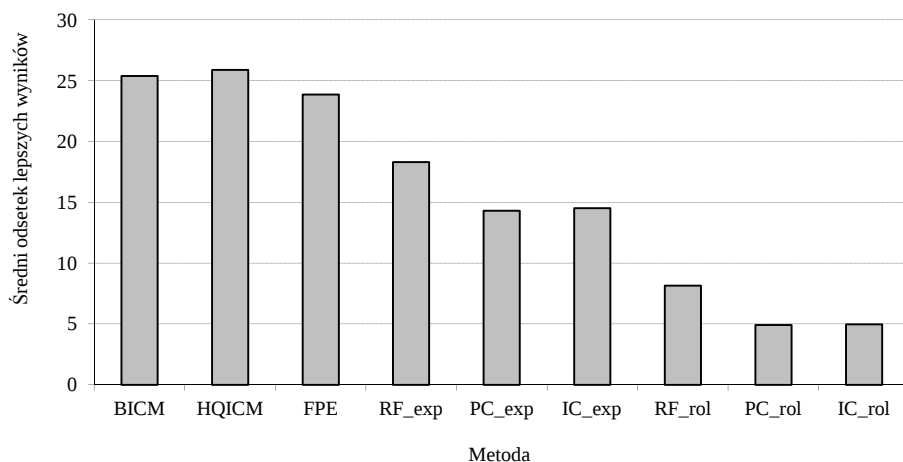
W badaniach symulacyjnych w zależności od liczby stawianych prognoz p , procedura generowania szeregów czasowych i testowania jakości prognoz powtarzana była 100 lub 1000 razy.

Własności różnych metod wyboru optymalnego modelu badano także na danych rzeczywistych. W tym celu wykorzystano szereg miesięcznej inflacji konsumenckiej w Polsce obejmujący okres 11.2000-08.2011. Liczył on więc 130 obserwacji. Zbiór zmiennych objaśniających składał się z 54 zmiennych, wśród których były m.in.: sektorowe indeksy cen, wskaźniki inflacji bazowej, indeksy produkcji sprzedanej przemysłu, produkcja ważniejszych wyrobów, zmienne reprezentujące dynamikę handlu zagranicznego, podstawowe zmienne rynku pracy, wyniki badań przewidywanej koniunktury gospodarczej, ceny surowców, poziom stóp procentowych oraz kursy walut. Wszystkie zmienne zostały pozbawione trendu oraz odsezonowane w programie Demetra. Tak przygotowane szeregi potraktowano tak, jak jeden z wysymulowanych szeregów opisanych powyżej, przy czym przy badaniu analizowano zarówno próby przesuwane, jak i rozszerzane. A więc na podstawie okien o długości co najmniej T_p wybierano najlepszy model według każdego kryterium, który następnie służył do obliczania prognoz. Podstawą oceny jakości prognoz były błędy średniokwadratowe oraz średnie błędy absolutne prognoz wygasłych oraz wyniki testów CPA.

4. Wyniki

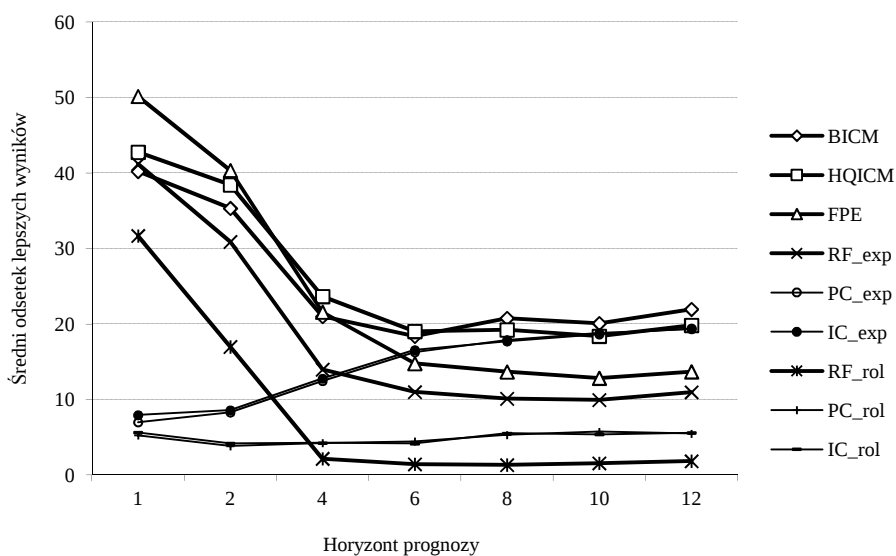
Dla danych symulowanych prezentowane są wyniki wskazujące, w jakim odsetku prób model wybrany określoną metodą okazał się lepszy z punktu widzenia bezwarunkowego testu CPA na poziomie istotności od innych modeli, przy czym dla zwiększenia czytelności prezentacji podawane są rezultaty uśrednione dla wszystkich konkurencyjnych modeli.

Wyniki uzyskano dla szeregów liczących 212 obserwacji, z czego szereg służący do wyboru najlepszego modelu liczył $T = 100$ obserwacji, a liczba przesuwanych okien była równa $p = 100$. Dokładności prognoz badano dla 7 różnych horyzontów: $h = 1, 2, 4, 6, 8, 10, 12$. Liczba zmiennych objaśniających została ustalona na poziomie $N = 100$. W przypadku metod wykorzystujących badanie dokładności prognoz wygasłych okna stosowane przy wyborze modelu liczyły $T_p = 50$ obserwacji. Liczba generowanych szeregów była równa $n = 1000$. We wszystkich przypadkach zbiór parametrów, względem którego dokonywano wyboru najlepszego modelu był postaci $K = \{1, 2, \dots, 6\}$, $n_y = \{0, 1\}$, $n_f = \{0, 1, 2\}$.



Rys. 1. Ranking metod – wariant podstawowy ($p = 100$, $n = 1000$, $\rho_F = 0,9$, $\rho_Z = 0$, stałe $\mathbf{L}$)

W podstawowym wariancie zakładano, że szereg czynników cechuje się autokorelacją na poziomie $\rho_F = 0,9$ (zob. równanie 11), wszystkie składniki losowe są homoskedastyczne, a macierz ładunków czynnikowych $\mathbf{L}$ jest stała w czasie. Wyniki dla takiego wariantu zilustrowano na rys. 1. Przedstawia on średnią liczbę przypadków, w których dana metoda okazała się lepsza od metod konkurencyjnych, przy czym podane są wyniki średnie ze względu na metody konkurencyjne oraz rozważane horyzonty prognozy. Z rysunku wynika więc, że metody BICM i HQICM okazywały się średnio w 25% spośród 1000 przypadków statystycznie lepsze od innych metod. Niewiele niższy wynik uzyskuje kryterium FPE. Metody wykorzystujące analizę jakości prognoz wygasłych na podstawie okien rozszerzanych uzyskują już wyniki gorsze – około 17% dla metody jednoczesnej RF_ext oraz poniżej 15% dla metod łączonych PC_exp i IC_exp. Najgorsze rezultaty uzyskują metody korzystające z okien przesuwanych. Dla podejść łączonych jest to mniej więcej 5%.

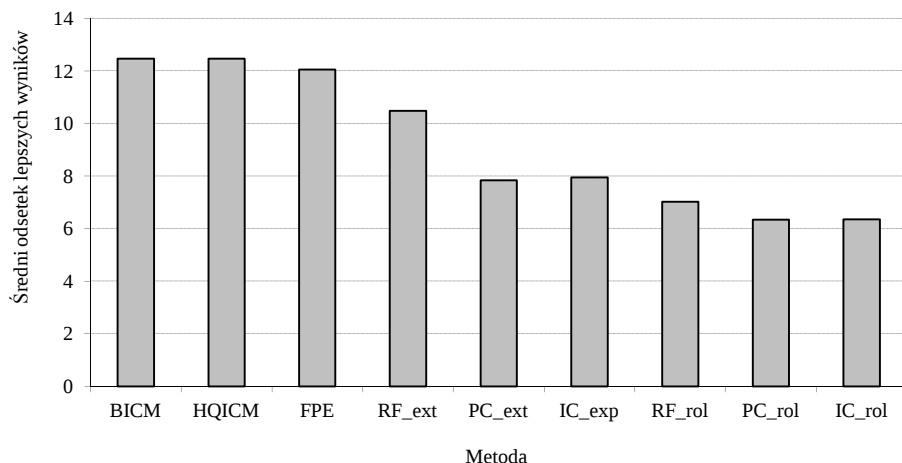


Rys. 2. Ranking metod według horyzontu prognoz – wariant podstawowy ($p = 100$, $n = 1000$, $\rho_F = 0,9$, $\rho_Z = 0$, stałe $\mathbf{L}$)

Rysunek 2 prezentuje bardziej szczegółowe wyniki z rozbiciem na różne horyzonty prognoz. Wynika z niego, że dla horyzontu $h = 1$ metody jednoczesne BICM, HQICM, FPE, RF_ext oraz RF_rol osiągają znacznie lepsze rezultaty od metod kombinowanych. Przykładowo kryterium FPE dawało lepsze wyniki od

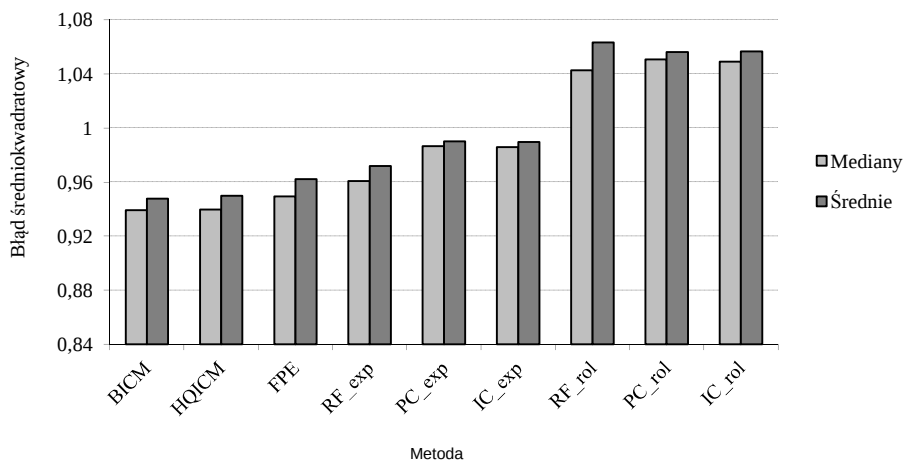
innych modeli średnio w 50% przypadków. Kryterium RF_rol było lepsze mniej więcej w 30% przypadków, a dla pozostałych było to w granicach 40%. Kryteria łączone dla tego horyzontu dawały lepsze wyniki w mniej niż 10% analizowanych przypadków, jednak w miarę wzrostu horyzontu przewaga metod jednoczesnych wyraźnie spada. Dla horyzontu $h = 12$ metody PC_exp oraz IC_exp dają praktycznie takie same rezultaty (około 20%), jak najlepsze kryteria jednoczesne BICM oraz HQICM. Zdecydowanie najsłabiej prezentuje się podejście RF_rol. Na omawianym rysunku widać także wyraźnie, że oba podejścia łączone wykorzystujące okna rozszerzane (PC_exp i IC_exp) dają praktycznie identyczne rezultaty. Podobne spostrzeżenie dotyczy metod PC_rol i IC_rol. Taka sytuacja ma miejsce we wszystkich wynikach prezentowanych w pracy.

Rysunek 3 ilustruje zbiorcze wyniki w przypadku stosowania wersji warunkowej testu CPA, przyjmując jako warunek ostatnią dostępną obserwację zmiennej prognozowanej: $S_i = Y_{T+i-1}$. Z takiej perspektywy różnice między podejściami są dużo mniejsze, choć pozycja metod w rankingu nie zmieniła się istotnie. Najlepsze kryteria jednoczesne okazywały się lepsze od innych metod średnio w ponad 12% przypadków. Najsłabsze wyniki osiągały podejścia wykorzystujące okna przesuwane – niewiele ponad 6%.



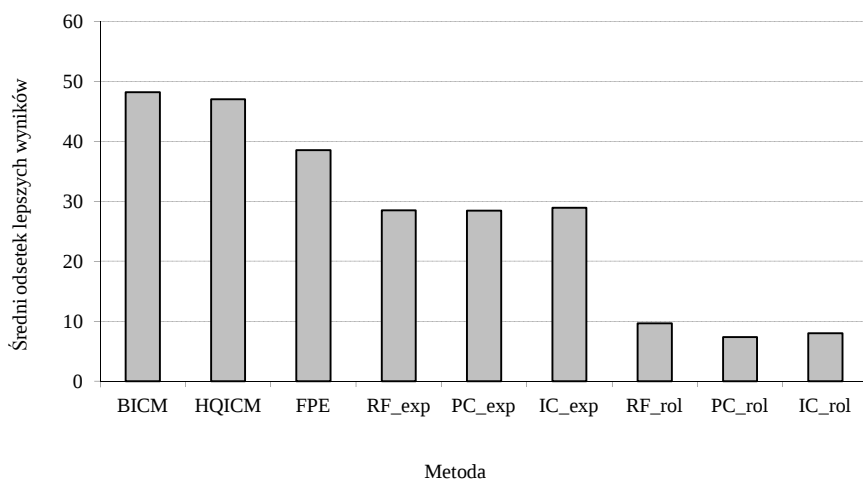
Rys. 3. Ranking metod – testy warunkowe, wariant podstawowy ($p = 100$, $n = 1000$, $\rho_F = 0,9$, $\rho_Z = 0$, stałe L)

Na rys. 4 przedstawiono charakterystyki błędów średniokwadratowych prognoz wygasłych – średnie wartości median i średnich dla różnych horyzontów. Wyniki te są dość zgodne z rankingiem przedstawionym na rys. 1. Najniższe błędy uzyskiwano dla jednoczesnych kryteriów informacyjnych, nieco wyższe w przypadku metod korzystających z okien rozszerzanych. Zdecydowanie najwyższymi błędami cechowały się podejścia związane z oknami przesuwanymi. Należy też zwrócić uwagę, że relatywne różnice pomiędzy błędami nie są duże. Rozpiętość wyników nie przekracza 10%.

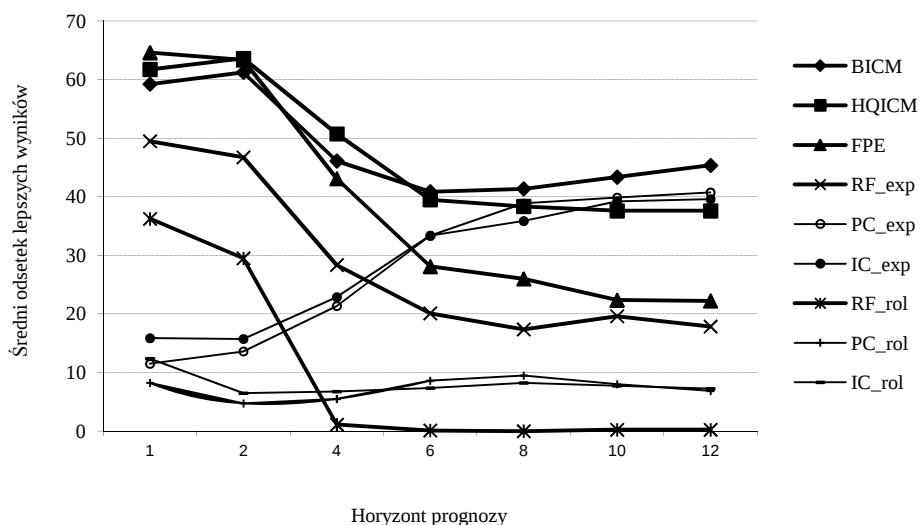


Rys. 4. Błędy średniokwadratowe dla wariantu podstawowego ($p = 100$, $n = 1000$, $\rho_F = 0,9$, $\rho_Z = 0$, stałe L)

Dla zweryfikowania rezultatów przedstawionych na rys. 1 i 2 przeprowadzono podobne badanie, ale przy większej liczbie prognoz $p = 500$ stawianej dla każdego symulowanego szeregu. Większa liczba obserwacji sprawia, że wyniki testów CPA są bardziej wiarygodne i częściej wskazują na odrzucenie hipotezy o braku różnic w jakości prognoz dla dwóch modeli, jednak ze względu na czasochłonność obliczeń dla takiej liczby prognoz przeprowadzono tylko 100 powtórzeń procedury. Wyniki badań prezentują rys. 5 i 6. Potwierdzają one rezultaty zaprezentowane dla wariantu podstawowego na rys. 1 i 2. Najlepsze wyniki dają kryteria BICM oraz HQICM. Nieco gorsze FPE i metody korzystające z prognoz wygasłych i okien rozszerzanych. Zdecydowanie najslabiej wypadają metody wykorzystujące okna przesuwane. Również rozłożenie wyników w zależności od horyzontu prognoz jest podobne jak na rys. 2.



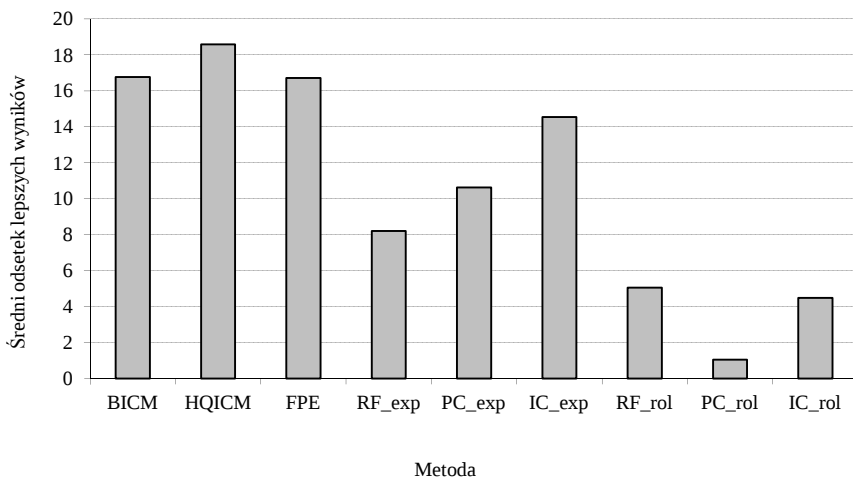
Rys. 5. Ranking metod – wariant podstawowy, większa liczba prognoz ($p = 500$, $n = 100$, $\rho_F = 0,9$, $\rho_Z = 0$, stałe L)



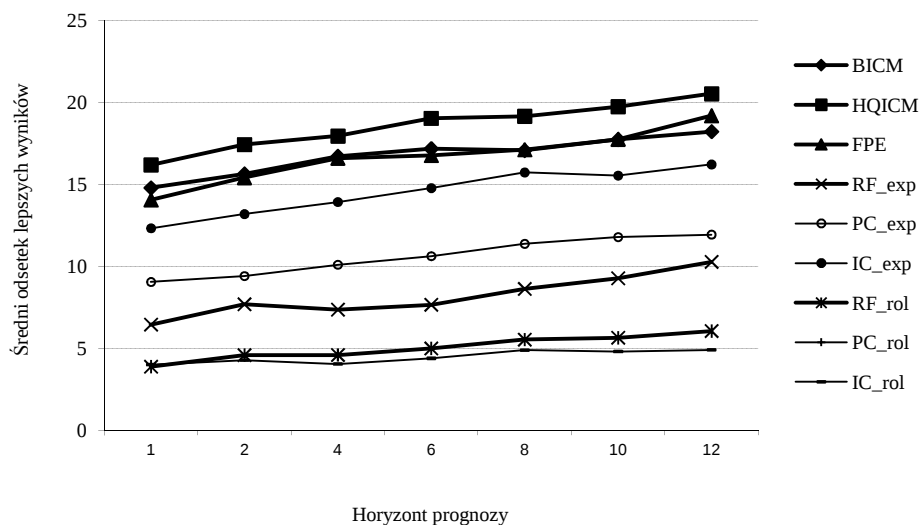
Rys. 6. Ranking metod według horyzontu prognoz – wariant podstawowy, większa liczba prognoz ($p = 500$, $n = 100$, $\rho_F = 0,9$, $\rho_Z = 0$, stałe L)

Należy zwrócić uwagę, że przy dużej liczbie prognoz testy zdecydowanie częściej wskazują na odrzucenie hipotezy o równych zdolnościach prognostycznych. W przypadku najlepszych kryteriów średnio w prawie 50% prób okazywały się one statystycznie lepsze od innych podejść, a dla krótkich horyzontów odsetek ten przekraczał 60%. Jest to więc mniej więcej dwa razy więcej niż w wariancie podstawowym.

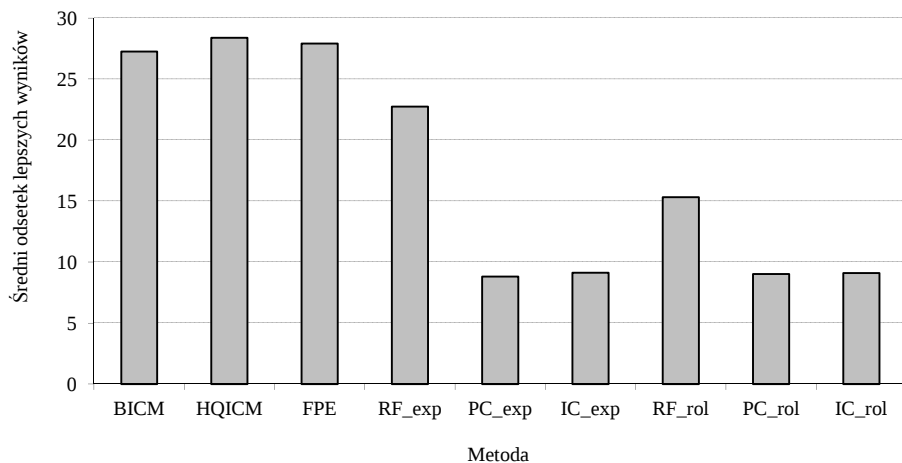
Rysunki 7 i 8 prezentują wyniki dla modelu czysto losowego, w którym nie występowała żadna autokorelacja $\rho_F = 0$, $\rho_Z = 0$. Także w tym przypadku najlepsze wyniki dają kryteria BICM, HQICM oraz FPE. Dobrze wypada także metoda IC_exp. Podobnie jak w poprzednich przypadkach, najgorsze prognozy otrzymuje się z modeli bazujących na prognozach wygasłych o oknach przesuwanych. W przeciwieństwie do poprzednich przypadków ranking ten jest stabilny ze względu na horyzont prognozy i prawie zawsze wygląda w taki sam sposób, co ilustruje rys. 8.



Rys. 7. Ranking metod – model czysto losowy ($p = 100$, $n = 1000$, $\rho_F = 0$, $\rho_Z = 0$, stałe L)



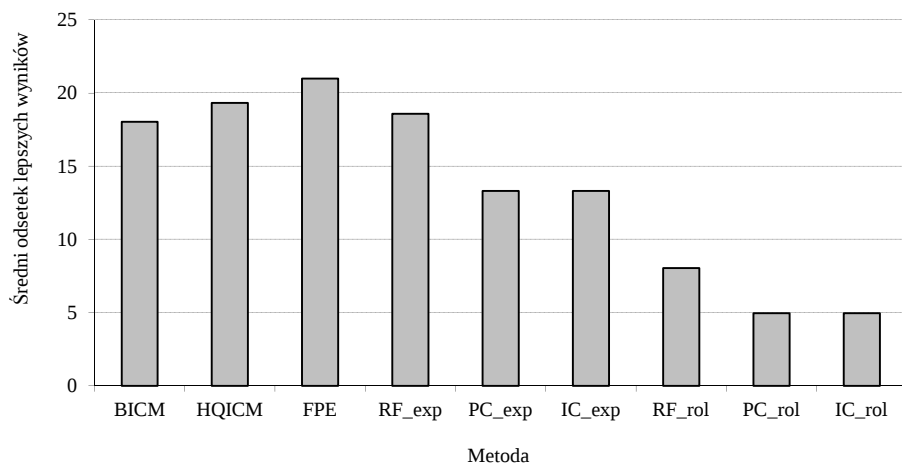
Rys. 8. Ranking metod według horyzontu prognoz – model czysto losowy ($p = 100$, $n = 1000$, $\rho_F = 0$, $\rho_Z = 0$, stałe L)



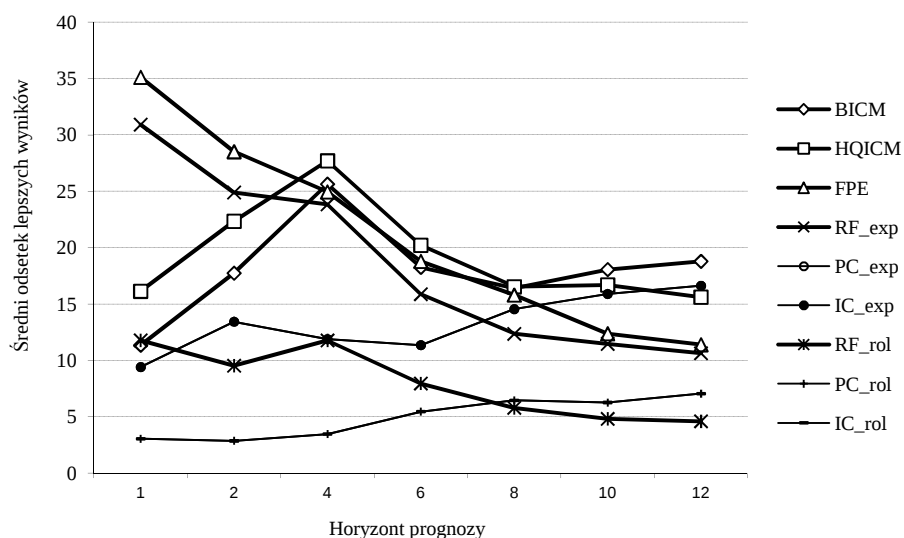
Rys. 9. Ranking metod – błędzenie losowe dla czynników ($p = 100$, $n = 1000$, $\rho_F = 1$, $\rho_Z = 0$, stałe L)

Na rys. 9 przedstawiono ranking metod w sytuacji, gdy dynamika czynników opisana jest procesem błędzenia losowego. Wyniki zbliżone są do uzyskanych w wariancie podstawowym, przy czym teraz najsłabiej wypadają kryteria mieszane, zarówno korzystające z okien przesuwanych, jak i rozszerzanych. Były one lepsze od innych podejść średnio w mniej niż 10% przypadków. Wyniki z rozbićm na różne horyzonty nie są prezentowane, gdyż są one zbliżone do tych z rys. 2.

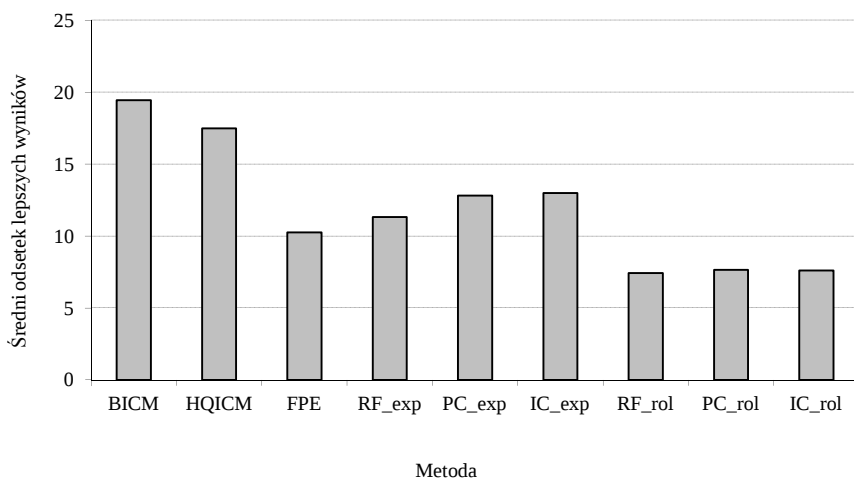
Rysunki 10 i 11 dotyczą modeli, w których zamiast autokorelacji czynników występuje autokorelacja składników losowych w równaniu (13). Należy zwrócić uwagę, że w takiej sytuacji wyniki różnią się w porównaniu do wariantu podstawowego. Najlepsze rezultaty przynosi stosowanie kryterium FPE (ponad 20%). Nieco mniej dają kryteria BICM, HQICM oraz RF_exp. Kryteria, w których stosowane są okna przesuwane dają najgorsze rezultaty. Dokładniejszy wgląd w wyniki daje rys. 11. Widać, że kryteria FPE oraz RF_exp mają zdecydowaną przewagę dla horyzontu $h = 1$. Dla $h = 4$ podobne wyniki dają również kryteria BICM oraz HQICM. Podobnie jak w wariancie podstawowym, dla dłuższych horyzontów poprawia się relatywna jakość prognoz dla modeli korzystających z rozszerzanych okien.



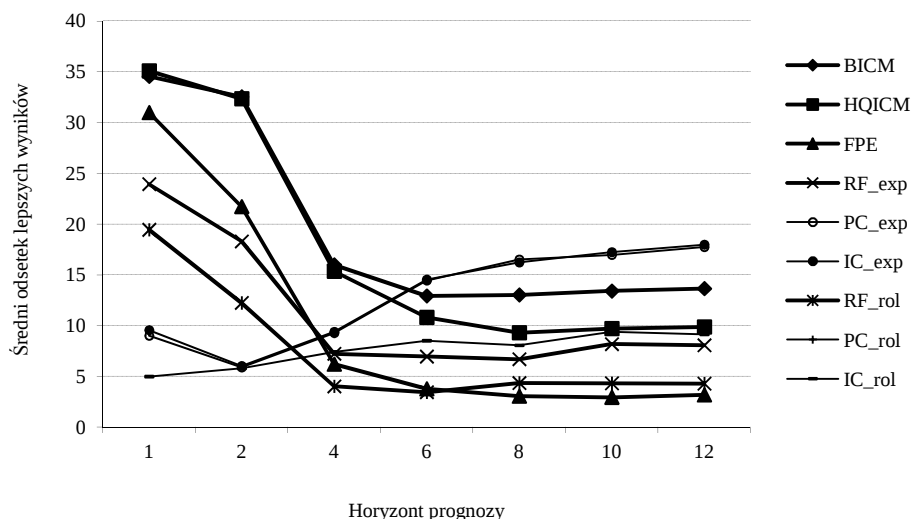
Rys. 10. Ranking metod – autokorelacja zaburzeń losowych ($p = 100$, $n = 1000$, $\rho_F = 0$, $\rho_Z = 0,9$, stałe L)



Rys. 11. Ranking metod według horyzontu prognoz – autokorelacja zaburzeń losowych ($p = 100$, $n = 1000$, $\rho_F = 0$, $\rho_Z = 0,9$, stałe L)



Rys. 12. Ranking metod – zmienna macierz ładunków czynnikowych ($p = 100$, $n = 1000$, $\rho_F = 0,9$, $\rho_Z = 0$, zmienne L)



Rys. 13. Ranking metod według horyzontu prognoz – zmienna macierz ładunków czynnikowych ($p = 100$, $n = 1000$, $\rho_F = 0,9$, $\rho_Z = 0$, zmienne L)

Na rys. 12 i 13 zaprezentowano wyniki dla zmieniającej się macierzy ładunków czynnikowych zgodnie z równaniem (12), w którym przyjęto $\psi_L = 0,015$ oraz $\rho_L = 0,9$. Widać, że nawet w sytuacji niestabilności procesu generującego dane, metody wykorzystujące analizę prognoz wygasłych z oknami przesuwanymi dają najgorsze wyniki. Ogólnie, rezultaty w takim ujęciu w dalszym ciągu są podobne do wyników uzyskanych dla wariantu podstawowego.

W dalszej części badano własności omawianych metod w odniesieniu do rzeczywistego szeregu inflacji konsumenckiej w Polsce. Szereg ten liczył 130 obserwacji. Na ich podstawie generowano $p = 48$ szeregów czasowych liczących $T = 70$ obserwacji lub więcej w przypadku okien rozszerzanych. Na ich podstawie wybierano najlepsze modele i obliczano wartości prognoz wygasłych oceniając ich dokładność. Dla metody wykorzystujących analizę prognoz wygasłych podokna miały długość co najmniej $T_p = 40$ obserwacji. Wyniki badań zestawiono w tab. 1-4, przy czym nie podawano rezultatów testów CPA, ponieważ w przeważającej liczbie przypadków nie dawały one podstaw do rozróżniania między zdolnościami prognostycznymi analizowanych podejść.

Tabela 1 zawiera błędy średniokwadratowe uzyskane dla prognoz przesuwanych w omawianym podejściu. Dwa najlepsze wyniki dla każdego horyzontu pogrubiono. Dwa najgorsze rezultaty oznaczono kolorem szarym. Podano także wyniki średnie, a w ostatniej kolumnie ranking metod ze względu na

średnie rezultaty. Najmniejsze błędy dawało stosowanie kryteriów BICM oraz HQICM, dobre wyniki także kryterium PC_exp. Najslabiej wypadły kryteria stosujące tylko analizę prognoz wygasłych RF_exp i RF_rol. Widać też dość wyraźnie, że kryteria łączone wykorzystujące okna przesuwane dawały gorsze wyniki niż te, które bazują na oknach rozszerzanych. Należy również podkreślić, że choć kryterium RF_exp daje najlepsze prognozy ze wszystkich metod dla horyzontu $h = 1$, to dla większych horyzontów daje często najgorsze wyniki.

Tabela 1

Błędy średniokwadratowe prognoz inflacji dla prognoz przesuwanych

Model	$h = 1$	$h = 2$	$h = 4$	$h = 6$	$h = 8$	$h = 10$	$h = 12$	Średnio	Rank
BICM	0,236	0,228	0,234	0,246	0,257	0,273	0,268	0,249	1
HQICM	0,234	0,234	0,235	0,246	0,257	0,272	0,268	0,249	1
FPE	0,232	0,261	0,237	0,246	0,255	0,276	0,278	0,255	5
RF_exp	0,226	0,271	0,245	0,255	0,293	0,288	0,288	0,267	8
PC_exp	0,232	0,239	0,235	0,249	0,256	0,272	0,267	0,250	3
IC_exp	0,239	0,239	0,236	0,249	0,266	0,273	0,271	0,253	4
RF_rol	0,238	0,268	0,238	0,240	0,324	0,295	0,363	0,281	9
PC_rol	0,242	0,248	0,237	0,242	0,264	0,294	0,290	0,260	6
IC_rol	0,248	0,247	0,239	0,242	0,268	0,298	0,288	0,261	7

Dosyć podobne rezultaty uzyskuje się przy zastosowaniu okien rozszerzanych przy wyznaczaniu kolejnych prognoz, co ilustruje tab. 2. Choć względne różnice między metodami są mniejsze niż w poprzednim przypadku, to wyraźniejszy jest ranking metod. Najlepsze są kryteria BICM, HQICM oraz FPE, nieco gorzej wypadają kryteria korzystające z okien rozszerzanych. Ponownie ostatnie pozycje zajmują metody stosujące okna przesuwane.

Tabela 2

Błędy średniokwadratowe prognoz inflacji dla prognoz rozszerzanych

Model	$h = 1$	$h = 2$	$h = 4$	$h = 6$	$h = 8$	$h = 10$	$h = 12$	Średnio	Rank
BICM	0,233	0,235	0,237	0,244	0,241	0,255	0,248	0,242	1
HQICM	0,237	0,240	0,237	0,244	0,241	0,255	0,248	0,243	2
FPE	0,218	0,263	0,237	0,242	0,241	0,255	0,244	0,243	2
RF_exp	0,227	0,277	0,238	0,247	0,247	0,254	0,247	0,248	6
PC_exp	0,237	0,241	0,237	0,245	0,248	0,260	0,247	0,245	4
IC_exp	0,237	0,241	0,237	0,245	0,248	0,260	0,247	0,245	4
RF_rol	0,270	0,287	0,240	0,258	0,298	0,287	0,276	0,274	9
PC_rol	0,245	0,251	0,243	0,253	0,258	0,274	0,280	0,258	7
IC_rol	0,245	0,251	0,243	0,253	0,258	0,274	0,280	0,258	7

W tab. 3 i 4 przedstawiono średnie błędy absolutne prognoz wygasłych. Z tego punktu widzenia najlepsze wyniki dawały kryteria BICM oraz HQICM, a także PC_exp oraz IC_exp. Ponownie zdecydowanie najslabsze wyniki były związane ze stosowaniem podejścia korzystającego z okien przesuwanych. Wyniki były więc ogólnie podobne do rezultatów uzyskanych przy rozważaniu błędów średniokwadratowych.

Tabela 3

Średnie błędy absolutne prognoz inflacji dla prognoz przesuwanych

Model	$h = 1$	$h = 2$	$h = 4$	$h = 6$	$h = 8$	$h = 10$	$h = 12$	Średnio	Rank
BICM	0,170	0,172	0,177	0,187	0,192	0,206	0,203	0,187	1
QHICM	0,173	0,175	0,178	0,188	0,192	0,206	0,203	0,188	3
FPE	0,181	0,203	0,179	0,189	0,190	0,211	0,208	0,194	5
RF_exp	0,172	0,209	0,185	0,193	0,218	0,222	0,219	0,203	6
PC_exp	0,164	0,181	0,177	0,189	0,191	0,203	0,202	0,187	1
IC_exp	0,176	0,180	0,179	0,189	0,199	0,203	0,209	0,191	4
RF_rol	0,185	0,210	0,184	0,187	0,236	0,241	0,290	0,219	9
PC_rol	0,180	0,190	0,182	0,190	0,203	0,242	0,232	0,203	6
IC_rol	0,192	0,189	0,184	0,189	0,208	0,243	0,231	0,205	8

Tabela 4

Średnie błędy absolutne prognoz inflacji dla okien rozszerzanych

Model	$h = 1$	$h = 2$	$h = 4$	$h = 6$	$h = 8$	$h = 10$	$h = 12$	Średnio	Rank
BICM	0,173	0,179	0,180	0,188	0,182	0,197	0,191	0,184	1
QHICM	0,176	0,182	0,180	0,188	0,182	0,197	0,191	0,185	2
FPE	0,164	0,212	0,180	0,184	0,182	0,197	0,188	0,187	5
RF_exp	0,183	0,225	0,181	0,191	0,189	0,195	0,190	0,193	6
PC_exp	0,176	0,185	0,180	0,188	0,187	0,199	0,190	0,186	3
IC_exp	0,176	0,185	0,180	0,188	0,187	0,199	0,190	0,186	3
RF_rol	0,203	0,229	0,184	0,194	0,227	0,233	0,223	0,213	9
PC_rol	0,184	0,188	0,184	0,193	0,201	0,222	0,227	0,200	7
IC_rol	0,184	0,188	0,184	0,193	0,201	0,222	0,227	0,200	7

Podsumowanie

Przedstawione w pracy wyniki w dość jednoznaczny sposób wskazują, że najlepsze modele DFM dla celów prognostycznych najczęściej były uzyskiwane na podstawie kryteriów jednoczesnych BICM oraz HQICM zaproponowanych przez Groena i Kapetaniosa. Tylko nieznacznie gorsze rezultaty były związane ze stosowaniem kryterium FPE Bai i Ng. Kolejne miejsca w rankingu przypadały kryteriom bazującym na oknach rozszerzanych przy wyborze najlepszego modelu. Natomiast najslabiej wypadały z reguły metody stosujące okna przesuwane. Spośród metod z dwóch ostatnich grup lepiej wypadały metody, w których proces wyboru najlepszego modelu przebiegał w całości w oparciu o badanie własności prognoz wygasłych. Wyniki wskazywały także, że w przypadku niewielkich horyzontów prognozy różnice pomiędzy metodami identyfikowane za pomocą testów zdolności prognostycznych CPA były z reguły znacznie wyraźniejsze niż w przypadku dłuższych horyzontów.

Literatura

1. Alessi L., Barigozzi M., Capasso., *A robust criterion for determining the number of factors in approximate factor models*, ECARES Working Paper 2009/023.
2. Artis M., Banerjee A., Marcellino M., *Factor forecasts for the UK*, „Journal of Forecasting” 2005, Vol. 24.
3. Bai J., Ng S., *Determining the number of factors in approximate factor models*, „Econometrica” 2002, Vol. 70 (1).
4. Bai J., Ng S., *Boosting diffusion indexes*, „Journal of Applied Econometrics” 2009, Vol. 24 (4).
5. Baranowski P., Leszczyńska A., Szafrński G., *Krótkookresowe prognozowanie inflacji z użyciem modeli czynnikowych*, „Bank i Kredyt” 2010, nr 41 (4).
6. Bernanke B., Boivin J., *Monetary policy in data-rich environment*, „Journal of Monetary Economics” 2003, Vol. 50 (3).
7. Breitung J., Eickmeier S., *Dynamic factor models*, Deutsche Bundesbank Discussion Paper Series 1: Economic Studies No 38/2005.
8. Doz C., Giannone D., Reichlin L., *A quasi maximum likelihood approach for large approximate dynamic factor models*, ECB Working Paper 674.
9. Forni M., Hallin M., Lippi M., Reichlin L., *The generalized dynamic factor model: Identification and estimation*, „The Review of Economics and Statistics” 2000, Vol. 82 (4).
10. Giacomini R., White H., *Tests of conditional predictive ability*, „Econometrica” 2006, Vol. 74 (6).
11. Groen J., Kapetanios G., *Model selection criteria for factor-augmented regressions*, Federal Reserve Bank of New York Staff Report No. 363, February 2009.

12. Marcellino M., Stock J., Watson M., *Macroeconomic forecasting in the Euro area: country specific versus Euro wide information*, „European Economic Review” 2003, Vol. 47.
13. Matheson T., *Factor model forecasts for New Zealand*, Reserve Bank of New Zealand DP2005/01.
14. Onatski A., *Testing hypotheses about the number of factors in large factor models*, „Econometrica” 2009, Vol. 77.
15. Otter P., Jacobs J., den Reijer A., *A criterion for the number of factors in a data-rich environment*, Mimeo 2011.
16. Stock J., Watson M., *Diffusion indexes*, NBER Working Paper 6702.

SELECTION CRITERIA FOR FORECASTING DYNAMIC FACTOR MODELS

Summary

The paper compares three groups of methods used for best dynamic factor model selection for forecasting: modified information criteria, methods exclusively based on ex post forecasts analysis and mixed algorithms. It searches for the approach that delivers best out-of-sample forecasts according to mean square error measure. The analysis utilizes both Monte Carlo generated samples as well as real time series used for forecasting consumer inflation in Poland. Results show that best forecasts are obtained from the modified information criteria proposed by Groen and Kapetanios, whereas the methods that employ ex post forecasts from rolling windows usually give the worst predictions.

Maciej Pichura

Uniwersytet Ekonomiczny w Katowicach

ANALIZA WPŁYWU KOSZTÓW TRANSAKCYJNYCH NA OCENĘ EFEKTYWNOŚCI WYBRANEJ STRATEGII INWESTYCYJNEJ

Wprowadzenie

Metody inwestowania kapitału na różnego rodzaju rynkach są przedmiotem badań i analiz praktycznie od momentu powstania regulacji organizujących i sankcjonujących ten obrót, choć miały one z pewnością miejsce również we wcześniejszych okresach. W wypadku giełd, czyli rynków kapitałowych o zorganizowanej formie obrotu, przedmiotem wykonywania transakcji są najczęściej instrumenty finansowe lub towary. Większość inwestorów uczestniczących w takim obrocie korzysta z usług pośredników – członków giełd. Członkowie giełdy mogą dokonywać na niej transakcji na własny rachunek lub na rachunek swoich klientów. Kiedy robią to na rachunek swoich klientów świadczą dla nich tzw. usługi brokerskie, stąd bardzo często nazywa się ich brokerami.

Zarówno giełda, która najczęściej jest spółką o odpowiednim statusie, jak i jej członkowie pobierają opłaty za wykonanie transakcji. W zamian za nie gwarantują jej prawidłowe rozliczenie oraz, w wypadku brokerów, reprezentują wszelkie interesy swoich klientów. Nie sposób zatem przy ustalaniu wyniku inwestycji na rynku kapitałowym nie uwzględnić kosztów w postaci opłat transakcyjnych. Nawet najprostsza inwestycja jest obciążona kosztem dwukrotnie, składa się bowiem z dwóch transakcji: kupna i sprzedaży. Rozwój badań nad inwestycjami kapitałowymi prowadzi do powstawania i modyfikacji metod inwestycyjnych, które nierzadko zakładają dokonywanie kilkudziesięciu, a nawet kilkuset transakcji. Do najpopularniejszych metod inwestycyjnych należą strategie budowania portfela inwestycyjnego, do których najczęściej dobiera się spółki na podstawie zasad analizy fundamentalnej, oraz strategie inwestycyjne

prorowadzone według zasad analizy technicznej. W metodach portfelowych można uwzględnić koszty transakcyjne w procesie doboru składników i udziałów portfela oraz ustalić te koszty dla przyszłych transakcji z dużą dokładnością. Strategie inwestycyjne tworzone według zasad analizy technicznej także pozwalają na dokładne ustalenie kosztów transakcyjnych podczas etapu ich konstruowania i ustalania ich parametrów, ale nie dają jednak możliwości dokładnego ich oszacowania w trakcie stosowania, ponieważ nie jest znana liczba transakcji dla strategii w danym przyszłym okresie.

1. Proces konstruowania strategii inwestycyjnych

Strategię inwestycyjną można zdefiniować jako zbiór fundamentalnych zasad obejmujący teorię i praktykę przygotowania oraz prowadzenia inwestycji, jej poszczególnych elementów i najważniejszych koncepcji. Strategia powinna określać taktykę działania na rynku oraz cele, do których te działania będą prowadzić. Wszystkie metody inwestycyjne stosowane do utworzenia strategii trzeba definiować w sposób obiektywny i jednoznaczny.

Strategie inwestycyjne konstruowane według zasad analizy technicznej bardzo często mają postać zautomatyzowaną, którą można porównać z algorytmem lub ciągiem algorytmów, czy też z modelem. Strategia inwestycyjna w takiej formie polega bowiem na ustaleniu algorytmów zajmowania i zwalniania pozycji na rynku w oparciu o metody analizy technicznej oraz bezwzględny stosowaniu tych algorytmów. W dalszych rozważaniach na ten temat termin „strategia inwestycyjna” będzie się odnosić do powyższej definicji.

Proces konstruowania i rozwijania algorytmicznych strategii inwestycyjnych najczęściej składa się elementów, które zostały przedstawione w tab. 1.

Tabela 1

Proces konstruowania strategii inwestycyjnych

Element	Opis
1	2
Sformułowanie koncepcji strategii	Ustalenie metody lub wskaźników analizy technicznej, które posłużą do generowania sygnałów zajęcia i opuszczenia pozycji na rynku
Specyfikacja algorytmów i reguł strategii	Ustalenie zasad zajmowania i opuszczania pozycji na rynku Ustalenie reguł zarządzania ryzykiem (zlecenia z tzw. poziomami <i>stop loss</i> , <i>take profit</i> oraz zlecenia typu <i>trailing stop</i> *) Ustalenie zasad zarządzania wielkością pozycji**

cd. tabeli 1

1	2
Wstępne testy na danych historycznych	Testy mają dać odpowiedź na kilka podstawowych pytań: czy strategia przynosi zyski (mogą one być nieznaczne); czy zyski ze strategii nie wynikają tylko z jednej lub kilku transakcji o bardzo wysokim zwrocie; czy zyski ze strategii nie wynikają ze zwrotów uzyskanych tylko w jednym okresie cyklu rynkowego?
Optymalizacja strategii	Polega na odnalezieniu takiego zestawu parametrów transakcji (co do zasady powinny one być mierzalne), który da w efekcie najwyższą efektywność zastosowania strategii. Do najczęściej używanych metooptymalizacji należy tzw. algorytm siły (<i>brute force algorythm</i>), który polega na obliczeniu wyniku strategii dla wszystkich możliwych kombinacji wartości jej parametrów. Na tym etapie konieczne jest uwzględnienie kosztów transakcyjnych
Ocena efektywności i stabilności uzyskanych wyników	Jako kryterium oceny może służyć miernik efektywności zastosowany do optymalizacji strategii. Istotny jest jednak szeroki horyzont oceny i zastosowanie dodatkowych narzędzi

* Zlecenie typu *stop-loss* (zatrzymaj stratę) pozwala na ustalenie poziomu ceny, przy którym ząjęta pozycja zostanie zamknięta. Jest to cena, dla której strata z transakcji przekracza maksymalną akceptowalną przez inwestora. Poziom tej ceny jest często ustalany jako pewien ułamek ceny bieżącej lub na podstawie zakresu zmienności cen z określonego okresu. Można go również ustalać na podstawie wartości wybranych wskaźników technicznych lub zidentyfikowanych formacji cenowych. Zlecenia z poziomem *take-profit* (weź zysk), który jest stałą wartością, po osiągnięciu której pozycja zostaje zamykana (analogicznie do zleceń *stop-loss*, jednak w tym wypadku zlecenie jest zamykane, gdy osiągnięty zostanie zysk satysfakcjonujący inwestora). Najdogodniejsze jest stosowanie dla każdej transakcji tzw. poziomu *trailing-stop* (postępujący stop), czyli stale modyfikowanego poziomu ceny wyjścia z transakcji. W momencie otwarcia transakcji poziom ten jest ustalany w podobny sposób, jak poziom zlecenia *stop-loss*, jednak w wypadku korzystnej zmiany ceny poziom ten modyfikowany jest w taki sposób, aby zapewnić możliwie maksymalny zysk z transakcji (lub zminimalizować stratę). W wypadku niekorzystnych zmian cen poziom *trailing-stop* nie ulega modyfikacji.

** Koncepcja ta jest często porównywana z koncepcją wysokości zakładu w grach losowych. Problem zarządzania wielkością pozycji jest skomplikowany głównie z powodu trudności w prawidłowym modelowaniu rozkładu cen na rynkach kapitałowych (co np. powoduje mniej trudności w grach losowych).

Źródło: Opracowanie własne na podstawie [8].

Kolejne etapy konstruowania strategii inwestycyjnych należą już w zasadzie do następnej fazy procesu inwestycyjnego – rzeczywistego zastosowania i rozwoju strategii. W związku z tym nie będą one tu przedmiotem rozważań. Analizie poddany zostanie natomiast w największym stopniu wpływ kosztów transakcyjnych na prawidłowe wnioski odnośnie do efektywności podczas etapu optymalizacji strategii. Weryfikacja tego wpływu zostanie przeprowadzona na danych spoza próby optymalizacyjnej.

Metody oceny efektywności strategii inwestycyjnych

Najprostszym sposobem oceny efektywności strategii inwestycyjnej jest porównanie wartości kapitału końcowego z jego wartością w momencie rozpoczęcia inwestycji, które jest definiowane jako stopa zwrotu. Stopa zwrotu oraz skumulowana stopa zwrotu (która opisuje zmiany wartości kapitału w czasie – tzw. krzywą kapitału) są bardzo przydatnymi narzędziami oceny rezultatów inwestycji i z pewnością nie można ich zaniedbywać. Efektywność strategii inwestycyjnej nie powinna być jednak analizowana jedynie w kontekście sprawnościowym, który jest odzwierciedlany właśnie w postaci zwrotu kapitału. Należy ją rozumieć jako próbę oceny, w jakim stopniu spełnione zostały oczekiwania inwestora odnośnie do przepływów kapitału, mając na uwadze poziom ryzyka, który akceptuje inwestor oraz stabilność wyników inwestycyjnych w dłuższym okresie.

Rezultaty inwestycji w postaci średniego lub krańcowego zwrotu kapitału coraz częściej są odnoszone do ryzyka ponoszonego, aby zwrot ten uzyskać. Ryzyko w przeważającej liczbie przypadków jest dla inwestycji kapitałowej wyrażane miarami zmienności. Do najbardziej powszechnych należą wariancja i odchylenie standardowe stóp zwrotu (a także, zdecydowanie rzadziej, wartości instrumentu finansowego). W wypadku aktywów kapitałowych ryzyko bardzo często jest jednak mierzone jako zmienność stóp zwrotu w kierunku ujemnym, a mniejszą wagę przywiązuje się do odchylen dodatnich. Takie postrzeganie ryzyka w obszarze instrumentów finansowych jest opisywane jako tzw. dolne ryzyko (DR – *downside risk*) [1]. Jednym z mierników efektywności inwestycji, który powstał dzięki zastosowaniu odpowiedniej miary DR jest współczynnik UPR (*upside potential ratio*), który oblicza się korzystając ze wzoru [7]:

$$UPR = \frac{\frac{1}{T} \sum_{i=1}^T \tau^+(r_i - E(r))}{\sqrt{\frac{1}{T} \sum_{i=1}^T \tau^-(r_i - E(r))^2}} \quad (1)$$

gdzie:

UPR – współczynnik UPR dla T stóp zwrotu z inwestycji,

r_i – stopa zwrotu z aktywu w i -tym okresie,

$E(r)$ – oczekiwana minimalna stopa zwrotu (często równa 0),

$\tau^+ = 1$ dla $r_{ip} > E(r)$, $\tau^+ = 0$ dla $r_{ip} \leq E(r)$,

$\tau^- = 1$ dla $r_{ip} \leq E(r)$, $\tau^- = 0$ dla $r_{ip} > E(r)$.

Jedną z najciekawszych miar efektywności inwestycyjnej, a jednocześnie prostych i niesprawiających trudności w obliczaniu jest miernik Omega. Informuje on o całkowitym wpływie postaci rozkładu stóp zwrotu na ocenę efektywności. Jego zastosowanie nie wymaga określania lub zakładania dokładnej postaci formalnej rozkładu, ale niezbędna jest wiedza na temat dystrybuanty empirycznej tego rozkładu. Współczynnik Omega przyjmuje ciągły rozkład stóp zwrotu, oblicza go się korzystając ze wzoru [4]:

$$\Omega(L) = \frac{\int_a^b [1 - F(r)] dr}{\int_a^L F(r) dr} \quad (2)$$

gdzie:

$\Omega(L)$ – wskaźnik Omega o progu L dla danego szeregu stóp zwrotu,

$F(r)$ – funkcja dystrybuanty rozkładu stóp zwrotu,

(a, b) – przedział wartości stóp zwrotu (w szczególnych przypadkach może mieć nieskończone granice),

L – progowa (wzorcowa) wartość stopy zwrotu uznawana za zyskową.

W przypadku przeprowadzania doświadczeń empirycznych zdecydowanie łatwiej stosować wzór na współczynnik Omega, który zakłada, że rozkład stóp zwrotu ma postać skokową:

$$\Omega(L) = \frac{\frac{1}{n} \sum_{i=1}^n \tau^+(r_i - L)}{\frac{1}{n} \sum_{i=1}^n \tau^-(L - r_i)} \quad (3)$$

gdzie:

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

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

n – liczebność obserwacji szeregu stóp zwrotu.

Pozostałe oznaczenia jak w (1) i (2).

Współczynnik Omega odzwierciedla wiele charakterystyk rozkładu stóp zwrotu, który poddawany jest analizie. Pozwala on również na takie dostosowanie do preferencji inwestora, które daje możliwość porównywania między

sobą szeregów stóp zwrotu. Jest to miernik, który w ostatnim dziesięcioleciu uzyskał duże zainteresowanie zarówno wśród inwestorów praktyków, jak i teoretyków zajmujących się badaniem przebiegu oraz efektów inwestycji kapitałowych. Współczynnik ten jest obliczany z zastosowaniem dystrybuanty empirycznej badanego rozkładu stóp zwrotu. Podobny tok postępowania może skłonić do analizy efektywności inwestycji poprzez wykonanie testu statystycznego weryfikującego czy badany rozkład stóp zwrotu jest zgodny z wzorcowym lub oczekiwanym. Można przyjąć, że wzorcowy rozkład stóp zwrotu uzyskiwany jest w procesie konstruowania strategii inwestycyjnej na etapie optymalizacji jej parametrów. Etap testowania strategii inwestycyjnej na danych spoza próby optymalizacyjnej lub etap jej bieżącego stosowania można przyjąć za okres tworzenia testowego (analizowanego) rozkładu stóp zwrotu. Mając dane dwa empiryczne szeregi stóp zwrotu można się podjąć weryfikacji statystycznej polegającej na sprawdzeniu czy te dwie próbki pochodzą z jednego rozkładu statystycznego. Jedną z powszechnie stosowanych w tym celu metod jest test dwóch próbek Kołmogorowa-Smirnowa. Statystyka tego testu skrótowo jest oznaczana jako statystyka K-S i ma następującą postać [5]:

$$D_{n,n'} = \sup_{-\infty < x < \infty} |F_{1,n}(x) - F_{2,n'}(x)| \quad (4)$$

gdzie:

$D_{n,n'}$ – statystyka K-S dla dwóch badanych prób,

$F_{1,n}(x)$ – dystrybuanta empiryczna dla próby wzorcowej,

$F_{2,n'}(x)$ – dystrybuanta empiryczna dla próby testowej.

Test Kołmogorowa-Smirnowa jest konstruowany przy zastosowaniu obszaru krytycznego rozkładu Kołmogorowa. Hipotezą zerową jest przypuszczenie, że próby pochodzą z jednego rozkładu i odrzuca się ją na poziomie istotności α , jeśli [5]:

$$\sqrt{\frac{nn'}{n+n'}} D_{n,n'} > K_\alpha \quad (5)$$

gdzie:

n – liczebność próby wzorcowej,

n' – liczebność próby testowej,

K – zmienna losowa o rozkładzie Kołmogorowa.

Pozostałe, oznaczenia jak we wzorze (4).

Opisana wyżej metoda weryfikacji statystycznej dotycząca dwóch prób empirycznych może mieć zastosowanie w ocenie efektywności strategii inwestycyjnych. Jeśli bowiem uzyskany w okresie optymalizacji wzorcowy szereg stóp zwrotu zostanie uznany przez inwestora za spełniający jego oczekiwania (czyli jest w określony sposób efektywny), to kolejny szereg otrzymany na danych spoza próby optymalizacyjnej powinien mieć taki sam rozkład, jak wzorcowy. Jeśli można statystycznie potwierdzić taki związek, to istnieją podstawy do twierdzenia, że skonstruowana strategia inwestycyjna jest efektywna.

Zastosowanie mierników UPR oraz Omega i testu K-S jako metod oceny efektywności strategii inwestycyjnych opisanych w kolejnej części opracowania posłuży do analizy efektywności strategii inwestycyjnych w kontekście kosztów transakcyjnych.

3. Proponowana strategia inwestycyjna

Przedstawiona wybrana strategia inwestycyjna została przeanalizowana w badaniu empirycznym. Strategia ta jest konstruowana według zasad omówionych wcześniej. Ma kilka cech, które w zasadzie mają postać parametrów wejściowych. Parametry te odnoszą się do elementów ograniczających ryzyko strategii. Cechą wyróżniającą opisywaną strategię od pozostałych opartych na metodach analizy technicznej są zasady otwierania i zamykania pozycji na rynku. Do ich sformułowania posłuży w głównej mierze jeden z najbardziej popularnych wskaźników analizy technicznej.

W przypadkach posługiwania się wskaźnikami technicznymi do formułowania zasad zajmowania pozycji rynkowej znajdują zastosowanie tzw. sygnały przecięcia. Uzyskuje się je poprzez obserwację krzywej wartości wskaźnika technicznego oraz krzywej (lub linii) wartości sygnału ustalonego dla tego wskaźnika. Rozróżnienie typów sygnału dokonywane jest na podstawie kierunku, z którego następuje przecięcie tych krzywych. Najczęściej w wypadku przecięcia krzywej (linii) sygnału przez krzywą wskaźnika od dołu generowany jest impuls do zajęcia pozycji długiej. Gdy przecięcie następuje od góry, generowany jest sygnał do otwarcia pozycji [3]. Sygnały te można w sposób bardziej sformalizowany opisać za pomocą przedstawionego poniżej algorytmu.

Sygnał zajęcia pozycji długiej (kupna) wyznaczany dla przecięcia krzywej lub linii sygnału przez krzywą wskaźnika technicznego oznaczono zmienną zero-jedynkową $L(ind, sig)$. Wartość 0 zmienna przyjmuje wtedy, gdy sygnał nie został wygenerowany, wartość 1 – gdy sygnał jest wygenerowany w momencie t . Warunki, aby spełnić równość $L(ind, sig) = 1$ są następujące:

$$\begin{cases} pos \leq 0 \\ ind_{t-2} < sig_{t-2} \\ ind_{t-1} \geq sig_{t-1} \\ ind_t > sig_t \end{cases} \quad (6)$$

gdzie:

${}_tL(ind, sig)$ – wartość zmiennej sygnału kupna w momencie t ,

pos – zmienna odzwierciedlająca zajmowaną na rynku pozycję w momencie t (0 – brak pozycji rynkowej, 1 – pozycja długa, -1 – pozycja krótka),

ind_t – wartość krzywej wskaźnika w momencie t ,

sig_t – wartość krzywej lub linii sygnału w momencie (w wypadku linii jest stała w czasie).

Analogicznie, sygnał otwarcia pozycji krótkiej (sprzedaży) dla przecięcia stałej linii sygnału przez krzywą określonego wskaźnika oznaczono zmienną zero-jedynkową ${}_tS(ind, sig)$. Zmienna ma wartość 0, gdy sygnał nie jest generowany oraz wartość 1, gdy występuje w momencie t . Warunki potrzebne do spełnienia równości ${}_tS(ind, sig) = 1$ to:

$$\begin{cases} pos \geq 0 \\ ind_{t-2} > sig_{t-2} \\ ind_{t-1} \leq sig_{t-1} \\ ind_t < sig_t \end{cases} \quad (7)$$

gdzie:

${}_tS(ind, sig)$ – wartość zmiennej sygnału sprzedaży w momencie t .

Pozostałe oznaczenia, jak we wzorze (6).

Wspólne czynniki większości strategii inwestycyjnych należą do grupy mechanizmów ograniczania ryzyka strategii inwestycyjnej. Są nimi poziomy *stop-loss* (symbol SL), *take-profit* (symbol TP) oraz *trailing-stop* (symbol TS). Podobnie jak parametry wynikające z wyboru reguł analizy technicznej do zajmowania pozycji na rynku, poziomy SL, TP i TS podlegają optymalizacji. Wartości tych parametrów będą określane w procencie ceny danego aktywu.

Dodatkowym wspólnym parametrem wszystkich strategii jest poziom ceny wejścia przy zajmowaniu pozycji (symbol EP – *entry price*). Dobieranie poziomu EP jest związane z zasadą tzw. potwierdzenia trendu. Jest on ustalany w ujęciu procentowym ceny i podobnie jak omawiane powyżej parametry strategii może być poddawany optymalizacji. Dla wygenerowanych sygnałów kupna cena wejściowa jest ustalana na poziomie przewyższającym cenę w momencie wygenerowania sygnału o wartość EP. Dla sygnałów sprzedaży cena wejściowa to cena w momencie wystąpienia sygnału pomniejszona o wartość EP. Jeśli zmiana trendu okaże się bardzo krótkotrwała, to cena wejściowa transakcji nie zostanie w ogóle osiągnięta, co doprowadzi do braku otwarcia (zamknięcia) pozycji na rynku. Jest to swoisty mechanizm chroniący strategię inwestycyjną przed wykonaniem zbyt dużej liczby transakcji, co powiększa koszty transakcyjne, oraz transakcji wynikających z błędnych sygnałów.

W zaproponowanej strategii inwestycyjnej jako wskaźnik, który generuje sygnały do zajęcia lub opuszczenia pozycji na rynku służy jeden z powszechnie stosowanych wskaźników analizy technicznej – wskaźnik zmian ROC (*rate of change*). Jest on uznawany za uniwersalne i skuteczne narzędzie analizy technicznej. ROC określa procentową zmianę bieżącej ceny aktywu w stosunku do jego ceny sprzed określonej liczby notowań. Oblicza się go korzystając ze wzoru [6]:

$${}_t ROC_{T1}(x) = \left(\frac{x_t}{x_{t-T1}} - 1 \right) \cdot 100 \quad (8)$$

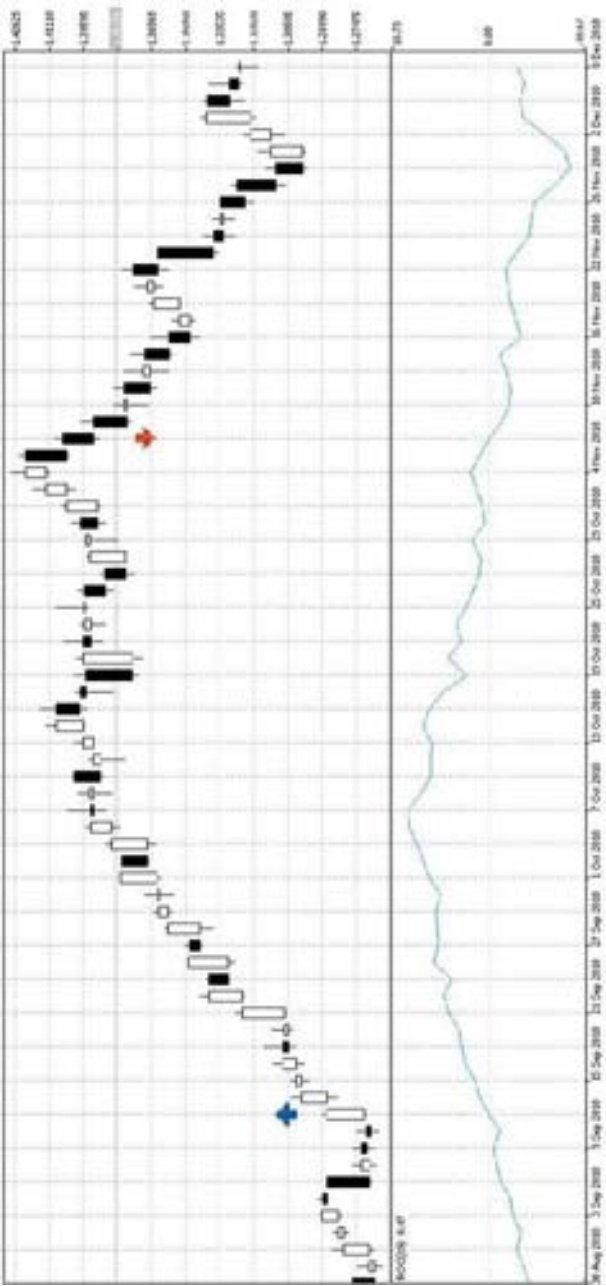
gdzie:

${}_t ROC_{T1}(x)$ – bieżąca wartość wskaźnika ROC dla aktywu x ,

$T1$ – okres zmian wskaźnika ROC,

x_t – bieżąca cena aktywu.

Wskaźnik ROC służy do sygnalizowania potencjalnych zmian kierunków trendów. Sygnały generowane przez ten wskaźnik są najczęściej analizowane na podstawie porównywania wartości jego krzywej do punktu centralnego, czyli linii zerowej. Jeśli krzywa ROC przecina linię zerową od dołu generowany jest sygnał kupna, jeśli przecięcie linii zerowej następuje od góry, wskazuje to na sygnał sprzedaży. Wykres wskaźnika ROC oraz przykładowe sygnały przez niego generowane są przedstawione na rys. 1.



Rys. 1. Wykres dzienny cen EUR/USD oraz wygenerowanego dla niego wskaźnika ROC_{19} za okres 30.08.2010-10.12.2010. Sygnał zajęcia pozycji długiej (kupna) oznaczono symbolicznie jako strzałkę w górę, a zajęcia pozycji krótkiej (sprzedaży) oznaczono symbolicznie jako strzałkę w dół

Strategia inwestycyjna utworzona z zastosowaniem wskaźnika ROC wykorzystuje sygnały do zajęcia pozycji opisane zmiennymi $L_t[{}_tROC_{Ti}(x), 0]$ oraz $S_t[{}_tROC_{Ti}(x), 0]$. Strategię tę można opisać następującymi parametrami:

$$STR_{ROC}(T1, SL, TP, TS, EP) \quad (9)$$

Najistotniejszym etapem po zdefiniowaniu podstawowych elementów strategii inwestycyjnej jest optymalizacja jej parametrów oraz weryfikacja efektywności na podstawie testów przeprowadzanych na okresie następującym po okresie optymalizacji. Podczas empirycznej weryfikacji efektywności przedstawionych strategii inwestycyjnych szczególna uwaga zostanie skierowana na aspekt wpływu kosztów transakcyjnych na ich rezultaty.

4. Analiza na podstawie badań empirycznych

Analizy empiryczne wykonane na potrzeby opracowania zostały przeprowadzone na danych historycznych uzyskanych z następujących źródeł: WIKIPOSIT, Yahoo Finance, BM BOŚ S.A. W celu uzyskania możliwie najbardziej różnorodnej informacji oraz możliwości wyciągnięcia wniosków dotyczących różnych rynków kapitałowych analizowane aktywa należą do różnych grup instrumentów. Są nimi: indeks WIG20, pary walutowe EUR/PLN i USD/PLN oraz kontrakty na złoto i ropę naftową. Dla każdego instrumentu wyróżniono okres optymalizacji testowanych strategii inwestycyjnych oraz okres weryfikacji ich efektywności. Okresy te starano się dobrać w taki sposób, aby zawierały przynajmniej jeden pełny cykl koniunkturalny, co pozwoli uniknąć błędów we wnioskowaniu, spowodowanych np. zbyt mało różnorodnym pod względem zmienności cen okresem inwestycyjnym. Lista wybranych do analizy instrumentów oraz zakres zastosowanych danych są przedstawione w tab. 2. Wszystkie dane zastosowane w badaniach mają charakter danych dziennych.

Tabela 2

Lista instrumentów kapitałowych wybranych do analizy oraz wyszczególnienie okresów optymalizacji i weryfikacji efektywności strategii

Nazwa instrumentu	Skrót nazwy	Okres optymalizacji	Okres weryfikacji
Indeks WIG20	WIG20	02.01.1994-31.12.2003	02.01.2004-01.10.2011
Para walutowa EUR/PLN	EUR/PLN	02.01.1999-31.12.2005	02.01.2006-01.10.2011
Para walutowa USD/PLN	USD/PLN	02.01.1994-31.12.2004	02.01.2005-01.10.2011
Kontrakt na ropę naftową CRUDE	OIL	02.01.1983-31.12.1999	02.01.2000-01.10.2011
Kontrakt na złoto	GOLD	02.01.1994-31.12.1996	02.01.1997-01.10.2011

Każdy przypadek symulacji inwestycji w dany instrument finansowy rozpatrywany jest na dwa sposoby. Pierwszy zakłada, że w okresie optymalizacji testowanej strategii inwestycyjnej nie są uwzględniane żadne koszty transakcyjne, a następnie wykonywana jest weryfikacja efektów tej strategii, w której koszty te występują (symulacja rzeczywistej inwestycji). Z drugiej strony analizie poddawana jest symulacja inwestycji poprzez zastosowanie wybranej strategii, która została zoptymalizowana z uwzględnieniem kosztów dokonywania transakcji (ustalonych na dwóch poziomach: 0,1% oraz 0,2% wartości transakcji). Jako kryteria optymalizacji strategii posłużyły wskaźniki: skumulowana stopa zwrotu oraz opisane wcześniej UPR i Omega, a także syntetyczny miernik powstały poprzez pomnożenie tych trzech wskaźników. Przedstawienie pełnych wyników optymalizacji oraz weryfikacji strategii z zastosowaniem założeń przyjętych w badaniu wiązałoby się ze znacznym powiększeniem objętości opracowania. W związku z tym zmniejszyłaby się jego czytelność, zatem opisano jedynie podsumowanie uzyskanych rezultatów oraz wnioski, które zostały wyciągnięte podczas ich analizy.

Rezultaty 80% optymalizacji, które uwzględniały koszty transakcji na poziomie 0,2% były różne od optymalizacji, w których nie uwzględniono kosztów transakcyjnych. W wypadku ustalenia kosztu transakcji w wysokości 0,1% parametry uzyskane po optymalizacji różniły się od parametrów uzyskanych przy założeniu braku takich kosztów dla 55% wykonanych optymalizacji. Informacje te pokazują, że koszty transakcyjne mają istotny wpływ na uzyskane podczas optymalizacji wartości parametrów strategii inwestycyjnej. Warte uwagi jest również przeanalizowanie, jaki wpływ na rezultaty inwestycyjne ma uwzględnienie kosztów transakcyjnych na etapie optymalizacji strategii.

Wyniki hipotetycznego zastosowania zaproponowanej strategii w okresie weryfikacji są zróżnicowane. W wypadku inwestycji w pary walutowe rezultaty są lepsze niż tzw. strategia „kup i trzymaj” i różnica ta jest znacząca (w wielu przypadkach co najmniej dwukrotna). Dla pary walutowej USD/PLN najlepsze efekty zarówno pod względem skumulowanej stopy zwrotu, jak i wskaźników Omega i UPR zostały uzyskane w wyniku optymalizacji strategii z uwzględnieniem najwyższego poziomu kosztów transakcyjnych. W wypadku pary walutowej EUR/PLN rezultaty końcowe inwestycji nie wykazują znaczących różnic, jeśli chodzi o zastosowany w okresie optymalizacji poziom kosztów transakcyjnych. Dla obydwóch wspomnianych par walutowych szereg stóp zwrotu uzyskany w okresie testowym pochodzi z tego samego rozkładu, co szereg uzyskany w okresie optymalizacji z uwzględnieniem kosztów transakcyjnych (poziom ufności testu K-S wynosi najczęściej powyżej 0,98). Brak uwzględnienia kosztów transakcyjnych na etapie optymalizacji skutkuje drastycznym obniżeniem

poziomu ufności tego testu szczególnie dla EUR/PLN. Dla USD/PLN w większości przypadków poziom ufności oscyluje na granicy wartości akceptowalnej (od 0,9 do 0,95).

Hipotetyczna inwestycja w indeks WIG20 dała zadowalający rezultat końcowy tylko w dwóch przypadkach na dwadzieścia cztery. Obydwa jednak wynikały z ustalenia podczas optymalizacji kosztów transakcyjnych na poziomie 0,2%. Zgodność statystyczna pod względem pochodzenia stóp zwrotu z jednego rozkładu zmiennej losowej została w wypadku inwestycji w WIG20 potwierdzona na wysokim poziomie ufności (powyżej 0,95) dla większości optymalizacji, które uwzględniały koszty transakcyjne zarówno na poziomie 0,2%, jak i 0,1%, natomiast znacznie rzadziej (tylko w dwóch z ośmiu przypadków) przy założeniu braku tych kosztów.

Najgorsze wyniki wykazały hipotetyczne inwestycje w kontrakty na ropę naftową i złoto – w zasadzie żadna z nich nie przyniosła zadowalających zysków. Być może było to spowodowane nieprawidłowo dobranym zakresem wartości optymalizowanych parametrów, jednak nie została przeprowadzona analiza przyczynowo-skutkowa w tym zakresie. Jeden z wyników może wydawać się zastanawiający – test K-S w ponad połowie badanych szeregów stóp zwrotu pozwala wnioskować na poziomie ufności ponad 0,95 o tym, że próbki te pochodzą z tego samego rozkładu co szeregi wzorcowe. Podobnie jak dla pozostałych analizowanych instrumentów, lepsze wyniki uzyskano przy uwzględnieniu kosztów transakcyjnych na etapie optymalizacji strategii.

Warto również zwrócić uwagę na fakt, że mierniki efektywności, których postać odzwierciedla kilka różnych cech rozkładu stóp zwrotu (UPR i Omega) zastosowane w strategii jako kryterium optymalizacji dają w rezultacie weryfikacji szeregi stóp zwrotu pochodzące z tych samych rozkładów (na podstawie testu K-S) co szeregi stóp zwrotu uzyskane podczas procesu optymalizacji.

Wnioski

Przeprowadzone badanie empiryczne nie dało rezultatów, które mogą prowadzić do w pełni jednoznacznych wniosków. Możliwe jest natomiast stwierdzenie, że kilka z nich pozwoliło na częściowe potwierdzenie hipotez stawianych jako cele weryfikacyjne opracowania.

Uwzględnienie kosztów transakcyjnych podczas optymalizacji parametrów strategii inwestycyjnych w większości analizowanych przypadków spowodowało uzyskanie innych zestawów parametrów strategii niż w wypadku, gdy kosztów tych nie uwzględniono. Dodatkowo test statystyczny K-S wykazał, że sze-

regi stóp zwrotu uzyskane w wyniku zastosowania strategii optymalizowanych z uwzględnieniem kosztów transakcyjnych pochodzą w zdecydowanie większej liczbie przypadków z tego samego rozkładu zmiennej losowej co szeregi wzorcowe (uzyskane w procesie optymalizacji) niż w sytuacji, gdy w optymalizacji założono brak kosztów transakcyjnych.

Zastosowanie jako kryteriów optymalizacji mierników, które odzwierciedlają kilka różnych cech rozkładu stóp zwrotu (Omega i UPR) również w większości przypadków dawało wynik testu K-S skutkujący brakiem podstaw do odrzucenia hipotezy zerowej. Zdecydowana większość tych testów wykonanych na szeregach stóp zwrotu uzyskanych poprzez optymalizację strategii na podstawie kryterium w postaci skumulowanej stopy zwrotu skłaniała do odrzucenia hipotezy zerowej.

Najbardziej zastanawiające i wymagające kolejnych analiz są końcowe wyniki inwestycyjne, które uzyskano poprzez zastosowanie zaproponowanej strategii. Najlepsze krańcowe stopy zwrotu uzyskano uwzględniając podczas optymalizacji koszty transakcyjne, jednak szeregi stóp zwrotu ponad połowy z najbardziej efektywnych strategii nie wykazały statystycznego podobieństwa do szeregów, które były wynikiem zoptymalizowanej strategii. Co więcej, tylko połowę rezultatów końcowych hipotetycznych inwestycji przeprowadzonych w badaniu można uznać za zadowalające.

Literatura

1. Ang A., Chen J.S., Xing Y., *Downside Risk*, „The Review of Financial Studies” 2006, Vol. 19, Issue 4.
2. Kaplan, P.D., Knowles J.A., Kappa A., *Generalized Downside Risk-Adjusted Performance Measure*, „Journal of Performance Measurement” 2004, Spring.
3. Katz J.O., McCormick D.L., *The Encyclopedia of Trading Strategies*, McGraw-Hill, New York 2000.
4. Keating C., Shadwick W.F., *A Universal Performance Measure*, „Journal of Performance Measurement” 2002, Vol. 6, No. 3.
5. Kryszewski W., Bartos J., Dyczka W., Królikowska K., Wasilewski M., *Rachunek prawdopodobieństwa i statystyka matematyczna w zadaniach. Część II. Statystyka matematyczna*, Wydawnictwo Naukowe PWN, Łódź 1999.
6. Le Beau C., Lucas D.W., *Komputerowa analiza rynków terminowych*, WIG Press, Warszawa 1998.
7. Le Sourd V., *Performance Measurement for Traditional Investment. Literature Survey*, EDHEC, Lille-Nice 2007.
8. Pardo R., *The Evaluation and Optimization of Trading Strategies*, Wiley & Sons, Hoboken, New Jersey 2008.

ANALYSIS OF INFLUENCE OF TRANSACTION COSTS ON INVESTMENT STRATEGIES EFFICIENCY

Summary

Main goal of this paper is to analyse the influence of transaction costs on efficiency evaluation and parameters selection regarding technical analysis investment strategies.

The first part of the article gives a brief description on rules of constructing automated investment strategies using technical analysis tools. Also, two efficiency measures that reflect several different features of distribution were presented: Omega and UPR (upside potential ratio). Moreover, a statistical method for testing whether a series of returns obtained in historical verification process comes from the same distribution as the series obtained during optimization process was introduced. Kolmogorov-Smirnov statistical test is used for that purpose.

Next, an investment strategy proposed for analysis was presented. A technical indicator ROC (rate of change) was used to generate signal for entering and exiting market positions. Afterwards, an empirical research was conducted on capital instruments of different markets: stock, currency and commodity.

The results of empirical research do not result in clear conclusions. Most of the analysed cases show that considering transaction costs at optimization level results in different strategy parameters selection than not doing so. However, efficiency of such strategies in many cases is not satisfactory despite statistical verification regarding the same distribution of given returns and returns of the most effective strategies in optimization process.

Agnieszka Przybylska-Mazur

Uniwersytet Ekonomiczny w Katowicach

REGUŁY POLITYKI PIENIĘŻNEJ A PROGNOZOWANIE WSKAŹNIKA INFLACJI

Wprowadzenie

Jednym z rodzajów polityki pieniężnej jest polityka oparta na regułach polityki pieniężnej. Ten rodzaj prowadzonej polityki umożliwia uruchamianie tzw. automatycznych stabilizatorów, dzięki którym gospodarka uzyskuje w polityce wsparcie dla zrównoważonego i stabilnego tempa wzrostu. Ponadto istnieje wówczas sprzężenie zwrotne pomiędzy instrumentami polityki pieniężnej a zmiennymi makroekonomicznymi, np. inflacją. Ze względu na stabilizujący charakter reguł polityki pieniężnej, w większości krajów dąży się do stosowania reguł odniesionych do inflacji jako celu polityki pieniężnej. W związku z tym decyzje dotyczące wysokości stóp procentowych mają wpływ na poziom inflacji. Problemem jest niedoskonałe sterowanie inflacją związane z opóźnieniem wpływu zmiany stóp procentowych na poziom inflacji, jak również z występowaniem szoków podażyowych i popytowych oraz z rodzajem stosowanego modelu. W związku z tym jednym z rozwiązań tego problemu jest rozważenie prognozy inflacji jako celu pośredniego w celu kontroli opóźnienia.

Jeżeli prognoza inflacji jest na poziomie celu inflacyjnego, to decydenci polityki pieniężnej nie powinni zmieniać stóp procentowych. Gdy powstaje zagrożenie zbyt wysoką inflacją przekraczającą wcześniej wytyczony cel, to decydenci polityki pieniężnej powinni podjąć decyzje o podniesieniu stóp procentowych w celu ograniczenia wzrostu cen. Natomiast jeżeli prognoza inflacji jest poniżej celu inflacyjnego, to powinna nastąpić obniżka stóp procentowych. Idea polega na niedopuszczeniu do przekroczenia wyznaczonego poziomu celu inflacyjnego. Taki sposób prowadzenia polityki nazywa się sterowaniem inflacją (*inflation targeting*).

Celem strategii bezpośredniego celu inflacyjnego realizowanej przez NBP jest utrzymanie inflacji na poziomie ustalonego celu inflacyjnego, jednak z powodu szoków podażyowych i popytowych niektóre odchylenia inflacji od celu in-

flacyjnego są nieuniknione. W związku z tym bank centralny powinien dążyć do minimalizacji odchyłeń inflacji od celu inflacyjnego. Ponadto, aby wyznaczyć prognozowaną wielkość inflacji w oparciu o wybrane modele strukturalne należy uwzględnić wysokość instrumentu polityki pieniężnej.

W opracowaniu zaprezentowano przykłady ilustrujące możliwość zastosowania reguł polityki pieniężnej do prognozowania wskaźnika inflacji.

1. Modele w prognozowaniu wskaźnika inflacji

Prognozy inflacji mogą być sporządzane różnymi metodami, ale wszystkie metody zapewniające ich przejrzystość i powtarzalność wymagają stosowania modeli.

Do prognozowania wskaźnika inflacji mogą być wykorzystane modele strukturalne. Istnieją modele strukturalne stanowiące klasę modeli dynamicznych, którą można zapisać w postaci przestrzeni stanów następująco [4]:

$$X_{t+1} = A \cdot X_t + B \cdot i_t + v_{t+1} \quad (1)$$

gdzie:

A – tzw. macierz towarzysząca,

B – wektor mnożników wpływu stóp procentowych,

X_t – wektor zmiennych stanu,

i_t – stopa procentowa,

v_{t+1} – wektor składników losowych.

W celu analitycznego wyznaczenia prognozy inflacji z uwzględnieniem wybranych reguł polityki pieniężnej w pracy został wykorzystany model Svenssona dla gospodarki narażonej na szoki podażowe i popytowe, który przedstawimy poniżej.

Przez π_t oznaczmy wskaźnik inflacji w okresie t , a przez π^* – cel inflacyjny.

Zgodnie ze strategią i założeniami, w ramach realizowanej przez NBP strategii bezpośredniego celu inflacyjnego od stycznia 2004 r. jest realizowany ciągły cel inflacyjny na poziomie 2,5% w ujęciu rok do roku, z symetrycznym przedziałem dopuszczalnych odchyłeń ± 1 punkt procentowy; realizacja ciągłego celu inflacyjnego oznacza, że odnosi się on do inflacji mierzonej w ujęciu miesiąc do analogicznego miesiąca poprzedniego roku.

Symbol i_t oznacza instrument polityki pieniężnej, np. stopę referencyjną, natomiast y_t względną luką pomiędzy rzeczywistym PKB Y_t i potencjalnym PKB Y_t^* wyrażoną w procentach, tzn. $y_t = 100 \cdot \frac{Y_t - Y_t^*}{Y_t^*}$.

Model strukturalny można wtedy zapisać następująco [2]:

$$\pi_{t+1} = \pi_t + \alpha y_t + \varepsilon_{t+1} \quad (2)$$

$$y_{t+1} = \beta_1 y_t - \beta_2 (i_t - \pi_t) + \eta_{t+1} \quad (3)$$

gdzie α , β_1 , β_2 są stałymi dodatnimi.

Składniki losowe ε_{t+1} , η_{t+1} mają rozkład o średniej równej zero, wariancjach równych σ_ε^2 , σ_η^2 i kowariancji $\sigma_{\varepsilon\eta}$. Składniki ε_{t+1} , η_{t+1} nie są obciążone autokorelacją. Składnik losowy ε_{t+1} przedstawia szok podażyowy, natomiast składnik losowy η_{t+1} szok popytowy.

Dla modelu opisanego równaniami (2), (3) mamy: wektor zmiennych stanu $X_t = \begin{bmatrix} \pi_t \\ y_t \end{bmatrix}$, macierz $A = \begin{bmatrix} 1 & \alpha \\ \beta_2 & \beta_1 \end{bmatrix}$. Wektory B i v_{t+1} są następujące: $B = \begin{bmatrix} 0 \\ -\beta_2 \end{bmatrix}$, $v_{t+1} = \begin{bmatrix} \varepsilon_{t+1} \\ \eta_{t+1} \end{bmatrix}$.

Powyższy model opisuje sytuację, w której zarówno inflacja, jak i zaregrowany popyt-produkcja reagują z opóźnieniem na zmianę instrumentu polityki pieniężnej. Z powyższego modelu wynika, że wzrost instrumentu banku centralnego powoduje spadek produkcji za jeden okres oraz spadek inflacji za dwa okresy. W związku z tym szczególnie ważna w podejmowaniu bieżących decyzji dotyczących wysokości instrumentu polityki pieniężnej, na podstawie modelu Svenssona gospodarki narażonej na szoki podażyowe i popytowe, jest prognoza inflacji na dwa okresy do przodu.

2. Reguły polityki pieniężnej

Przy prowadzeniu polityki opartej na regułach polityki pieniężnej bank centralny staje się źródłem pewności i stabilizacji oczekiwań. Polityka banku centralnego, a w szczególności reguły polityki monetarnej wywierają systematyczny wpływ na zmienne makroekonomiczne. Polityka oparta na regułach polityki pieniężnej może być aktywna lub pasywna. Gdy prowadzona jest aktywna

polityka pieniężna oparta na regułach, to istnieje sprzężenie zwrotne między rozważanymi zmiennymi makroekonomicznymi a instrumentami polityki pieniężnej. Ponieważ aktywne reguły polityki pieniężnej mają stabilizujący charakter, w większości krajów dąży się do stosowania tego rodzaju reguł, odniesionych do inflacji jako celu polityki pieniężnej.

Reguły polityki pieniężnej można podzielić na dwie klasy:

- reguły instrumentalne (*instrument rules*),
- reguły nastawione na cel (*targeting rules*).

Reguła instrumentalna wyraża instrument polityki pieniężnej jako funkcję dostępnej informacji o rzeczywistości.

Wyróżniamy następujące reguły instrumentalne:

- klasyczna reguła Friedmana dla bazy pieniężnej,
- standardowa reguła Taylora,
- modyfikacje standardowej reguły Taylora,
- reguła tylko inflacji,
- reguła luki bezrobocia.
- reguła Calvo.

2.1. Standardowa reguła Taylora

Według tej reguły stopy procentowe banku centralnego powinny zmieniać się bardziej niż odchylenie inflacji od pewnego ustalonego celu inflacyjnego, ale nie dąży się bezpośrednio do osiągnięcia tego celu, cel jest tylko jednym z argumentów funkcji decyzyjnej. Realizując standardową regułę Taylora bank centralny oprócz inflacji stabilizuje także produkcję.

Standardową regułę Taylora można zapisać w postaci [2]:

$$i_t = \pi^* + \alpha_z y_t + \beta_z (\pi_t - \pi^*) + r_t^* \quad (4)$$

gdzie:

- i_t – nominalna stopa procentowa banku centralnego (np. stopa referencyjna lub redyskontowa),
- π_t – poziom inflacji w okresie t ,
- π^* – cel inflacyjny,
- y_t – luka produkcyjna w okresie t ,
- r_t^* – realna stopa procentowa równowagi w okresie t .

W standardowej regule Taylora luka produkcji jest ważnym czynnikiem korygującym politykę pieniężną. Przy ustalonym współczynniku α_z można założyć, że im produkcja rzeczywista jest niższa od potencjalnej, tym bardziej powinna być obniżona stopa procentowa. Gdy produkcja obserwowana jest wyższa od długookresowego trendu, to zauważalne są wówczas symptomy przegrzania koniunktury, w efekcie czego ceny będą miały tendencję do wzrostu, natomiast decydenci polityki pieniężnej powinni wtedy podjąć decyzję o podwyższeniu stopy procentowej.

Zgodnie ze standardową regułą Taylora, w celu zmniejszenia inflacji należy zwiększyć realną stopę procentową, dzięki czemu zostanie ostudzona produkcja i tym samym zmniejszy się presja inflacyjna. W rozważanej regule stopa procentowa, czyli ogólnie instrument polityki pieniężnej zależy również od wartości współczynników reakcji. Z powodu braku jednoznaczności wyznaczenia instrumentu polityki pieniężnej w oparciu o standardową regułę Taylora, wynikającego z braku jednoznaczności wyznaczenia luki produkcyjnej, realnej stopy procentowej oraz współczynników reakcji, rozważa się również pewne modyfikacje standardowej reguły Taylora polegające na:

- pominięciu realnej stopy procentowej równowagi, wówczas stopę procentową wyznaczamy na podstawie wzoru:

$$i_t = \pi^* + \alpha_z y_t + \beta_z (\pi_t - \pi^*) \quad (5)$$

- usunięciu luki produkcyjnej:

$$i_t = \pi^* + \beta_z (\pi_t - \pi^*) + r_t^* \quad (6)$$

- zastąpieniu luki produkcyjnej wielkością produkcji:

$$i_t = \pi^* + \alpha_z Y_t + \beta_z (\pi_t - \pi^*) + r_t^* \quad (7)$$

gdzie:

Y_t – wielkość produkcji.

2.2. Reguła tylko inflacji

W regule tylko inflacji instrument polityki pieniężnej zależy tylko od inflacji oraz od celu inflacyjnego, co zapisujemy następująco:

$$i_t = \pi^* + \beta_z (\pi_t - \pi^*) \quad (8)$$

2.3. Reguła luki bezrobocia

Reguła luki bezrobocia ma postać:

$$i_t = \pi^* + \delta_z (u_t - u^*) + \beta_z (\pi_t - \pi^*) + r_t^* \quad (9)$$

gdzie:

u_t – aktualna stopa bezrobocia,

u^* – naturalna stopa bezrobocia.

Pozostałe symbole występujące we wzorze objaśniono wcześniej.

2.4. Reguła typu Calvo stopy procentowej

Zgodnie z koncepcją Calvo przedstawioną w 1983 r. zakłada się, że istnieje stałe prawdopodobieństwo kontynuowania reguły w każdym z okresów równe φ , natomiast prawdopodobieństwo zmiany reguły w każdym okresie wynosi $1 - \varphi$. Innymi słowy, w każdym okresie decydent podejmuje decyzję, czy zmienić wysokość instrumentu polityki pieniężnej czy pozostawić ją bez zmian.

Reguła typu Calvo stopy procentowej jest następującej postaci:

$$i_t = \begin{cases} h i_{t-1} + \beta_\pi (1-h) \theta_t & \text{dla } h \in [0,1), \beta_\pi > 0 \\ i_{t-1} + \gamma_\pi \theta_t & \text{dla } h=1, \gamma_\pi > 0 \end{cases} \quad (10)$$

Ta reguła opisuje sprzężenie zwrotne instrumentu polityki pieniężnej z prognozą inflacji, występujące w średnim horyzoncie czasowym $\frac{\varphi}{1-\varphi}$.

We wzorze (10) θ_t jest sumą dyskontowaną przyszłych oczekiwanych wskaźników inflacji:

$$\theta_t = (1-\varphi) E_t(\pi_t + \varphi \pi_{t+1} + \varphi^2 \pi_{t+2} + \dots) = (1-\varphi) \sum_{j=0}^{\infty} \varphi^j E(\pi_{t+j}), \quad (11)$$

$$\varphi \in (0,1)$$

Parametr h , $h \in [0,1]$ mierzy stopień wygładzenia stopy procentowej. Im większy parametr h , tym większy stopień wygładzenia.

Parametr β_π , $\beta_\pi > 0$ jest parametrem sprzężenia zwrotnego. Im większy jest parametr sprzężenia zwrotnego, tym szybciej eliminowana jest luka pomiędzy oczekiwaną inflacją i celem inflacyjnym.

3. Prognoza inflacji zgodna z regułą instrumentalną polityki pieniężnej

Prognoza wskaźnika inflacji $\pi_{t+T/t}(i_t)$ na T okresów do przodu wyznaczona w okresie t wyraża się wzorem:

$$\pi_{t+T/t}(i_t) = K_\pi X_{t+T/t}(i_t)$$

dla odpowiednio zdefiniowanego wektora K_π , zależnego od modelu, na podstawie którego wyznaczana jest prognoza wskaźnika inflacji.

Dla rozważanego w pracy modelu Svenssona $K_\pi = e_1 = [1 \ 0]$.

Wzór na prognozę wskaźnika inflacji $\pi_{t+T/t}(i_t)$, zgodną z regułą polityki pieniężnej obliczamy ze wzoru:

$$\pi_{t+T/t}(i_t) = K_\pi A^T X_t + K_\pi \sum_{j=1}^T A^{T-j} B i_{t+j-1/t} \quad (12)$$

Jeżeli prognoza wskaźnika inflacji wyznaczana jest na podstawie modelu Svenssona, to powyższy wzór przyjmuje postać:

$$\pi_{t+T/t}(i_t) = e_1 A^T X_t + e_1 \sum_{j=1}^T A^{T-j} B i_{t+j-1/t}$$

Wzory, na podstawie których wyznacza się prognozę wskaźnika inflacji w zależności od rodzaju stosowanej reguły polityki pieniężnej zestawiono w tab. 1.

Tabela 1

Zależność prognozy wskaźnika inflacji od rodzaju reguły polityki pieniężnej

Rodzaj reguły polityki pieniężnej	Wzory na prognozę wskaźnika inflacji na T okresów do przodu $\pi_{t+T/t}(i_t)$
1	2
Standardowa reguła Taylora	$K_\pi A^T X_t + K_\pi \sum_{j=1}^T A^{T-j} B$ $\cdot (\pi^* + \alpha_z y_{t+j-1/t} + \beta_z (\pi_{t+j-1} - \pi^*) + r_{t+j-1}^*)$

cd. tabeli 1

1	2
Modyfikacja standardowej reguły Taylora polegająca na zastąpieniu luki produkcyjnej wielkością produkcji	$K_{\pi} A^T X_t + K_{\pi} \sum_{j=1}^T A^{T-j} B \cdot (\pi^* + \alpha_z Y_{t+j-1/t} + \beta_z (\pi_{t+j-1} - \pi^*) + r_{t+j-1}^*)$
Modyfikacja standardowej reguły Taylora polegająca na pominięciu realnej stopy procentowej równowagi	$K_{\pi} A^T X_t + K_{\pi} \sum_{j=1}^T A^{T-j} B (\pi^* + \alpha_z y_{t+j-1/t} + \beta_z (\pi_{t+j-1/t} - \pi^*))$
Reguła tylko inflacji	$K_{\pi} A^T X_t + K_{\pi} \sum_{j=1}^T A^{T-j} B (\pi^* + \beta_z (\pi_{t+j-1/t} - \pi^*))$
Reguła luki bezrobocia	$K_{\pi} A^T X_t + K_{\pi} \sum_{j=1}^T A^{T-j} B \cdot (\pi^* + \delta_z (u_{t+j-1/t} - u_{t+j-1}^*) + \beta_z (\pi_{t+j-1/t} - \pi^*) + r_{t+j-1/t}^*)$
Reguła typu Calvo	$\begin{cases} K_{\pi} A^T X_t + K_{\pi} \sum_{j=1}^T A^{T-j} B (h i_{t+j-2/t} + \beta_{\pi} (1-h) \theta_{t+j-1/t}) & \text{dla } h \in [0,1), \beta_{\pi} > 0 \\ K_{\pi} A^T X_t + K_{\pi} \sum_{j=1}^T A^{T-j} B (i_{t+j-2/t} + \gamma_{\pi} \theta_{t+j-1/t}) & \text{dla } h=1, \gamma_{\pi} > 0 \end{cases}$

Teoretycznie rodzaj stosowanej reguły instrumentalnej ma wpływ na prognozę wskaźnika inflacji. Jeżeli prognoza wskaźnika inflacji jest wyznaczana na podstawie modelu Svenssona, to należy w powyższych wzorach przyjąć $K_{\pi} = e_1$.

4. Analiza empiryczna

Do analiz wzięto pod uwagę dane dotyczące miesięcznych wskaźników inflacji, instrumentu polityki pieniężnej – stopy referencyjnej oraz PKB. Ponieważ zgodnie ze strategią i założeniami, w ramach strategii bezpośredniego celu inflacyjnego od stycznia 2004 r. realizowany jest ciągły cel inflacyjny na poziomie 2,5% w ujęciu rok do roku, z symetrycznym przedziałem dopuszczalnych odchyłeń +/- 1 punkt procentowy oraz realizacja ciągłego celu inflacyjnego oznacza, że odnosi się on do inflacji mierzonej w ujęciu miesiąc do analogicznego miesiąca poprzedniego roku, a nie jak w latach 1999-2003, wyłącznie w grudniu do grudnia poprzedniego roku, dokonano analizy danych miesięcznych z okresu styczeń 2004 r.-marzec 2011 r. Oszacowane parametry modelu Svenssona wynoszą: $\alpha = -0,00619$, $\beta_1 = 0,3207$, $\beta_2 = 0,3901$.

W tab. 2 przedstawiono prognozy wskaźnika inflacji na jeden, dwa i trzy miesiące do przodu zgodne z regułą tylko inflacji.

Tabela 2

Prognozy wskaźnika inflacji zgodne z regułą tylko inflacji

β_z	Prognoza inflacji dla		
	T = 1 (kwiecień 2011 r.)	T = 2 (maj 2011 r.)	T = 3 (czerwiec 2011 r.)
0,01	4,26	4,24	4,23
0,1	4,26	4,24	4,24
0,5	4,26	4,25	4,24
1	4,26	4,25	4,24
1,5	4,26	4,25	4,25
5	4,26	4,27	4,28
10	4,26	4,29	4,33
100	4,26	4,68	5,34
650	4,26	7,07	15,15

Wyznaczając prognozę inflacji zgodną z regułą tylko inflacji możemy zauważyć, że inflacja w kolejnych trzech miesiącach nieznacznie będzie się obniżać, gdy współczynnik reakcji będzie przyjmował wartości 0,01; 0,1; 0,5; 1; 1,5. Natomiast im większa wartość współczynnika reakcji tendencja zmian wskaźnika inflacji jest przeciwna, wskaźnik inflacji będzie wzrastał nieznacznie dla $\beta_z = 5$. Im większa wartość współczynnika reakcji, tym większe tempo zmian wskaźnika inflacji.

W tab. 3 zestawiono prognozy wskaźnika inflacji na jeden, dwa i trzy miesiące do przodu zgodne z wybranymi regułami instrumentalnymi dla wybranych współczynników reakcji $\alpha_z = 0,5$, $\beta_z = 1,5$.

Tabela 3

Prognozy wskaźnika inflacji zgodne z wybranymi
regułami instrumentalnymi (dla $\alpha_z = 0,5$, $\beta_z = 1,5$)

Rodzaj reguły polityki pieniężnej	Prognoza inflacji dla		
	T = 1 (kwiecień 2011 r.)	T = 2 (maj 2011 r.)	T = 3 (czerwiec 2011 r.)
Standardowa reguła Taylora	4,26	4,27	4,28
Modyfikacja standardowej reguły Taylora polegająca na pominięciu realnej stopy procentowej równowagi	4,26	4,26	4,42
Modyfikacja standardowej reguły Taylora polegająca na zastąpieniu luki produkcyjnej wielkością produkcji	4,26	4,39	4,57
Reguła luki bezrobocia	4,26	4,26	4,27
Rzeczywista wartość wskaźnika inflacji	4,5	5	4,2

Na podstawie przeprowadzonej analizy prognoz wskaźnika inflacji zgodnych z wybranymi regułami instrumentalnymi, dla reguł instrumentalnych prezentowanych w tab. 3 można również wyciągnąć wnioski, że im większe wartości współczynników reakcji, tym wyższe wartości wyznaczonych prognoz wskaźnika inflacji.

Gdy wyznaczamy prognozę wskaźnika inflacji na podstawie modelu Svenssona zgodną z zaprezentowanymi regułami instrumentalnymi, to rodzaj zaprezentowanych reguł instrumentalnych nie ma znacznego wpływu na prognozę wskaźnika inflacji, ponieważ istnieje zależność pomiędzy kształtowaniem się

dynamiki produkcji i z nią związanej luki produkcyjnej oraz stopą bezrobocia. Ponadto wykorzystując do analiz model Svenssona dla horyzontu prognozy równego jeden rodzaj stosowanej reguły nie ma wpływu na wartość wskaźnika inflacji. Im większy horyzont prognozy, tym wpływ rodzaju stosowanej reguły instrumentalnej jest bardziej widoczny.

Zakończenie

Pomiędzy regułami polityki pieniężnej a prognozą wskaźnika inflacji istnieje sprzężenie zwrotne. Gdy prowadzona jest polityka pieniężna banku centralnego zwana sterowaniem inflacją, to bank centralny oszacowuje i ogłasza publicznie prognozowaną lub celową stopę inflacji i wtedy próbuje sterować aktualną inflacją w kierunku celu przez zmiany stopy procentowej i innych narzędzi kształtujących inflację. Można również wyznaczyć prognozy wskaźnika inflacji zgodne z regułami instrumentalnymi.

Literatura

1. Rudebush G.D., Svensson L.E.O., *Policy rules for inflation targeting*, „Working Paper Series”, National Bureau of Economic Research Cambridge 1998.
2. Svensson L.E.O., *Commentary: How Should Monetary Policy Respond to Shocks While Maintaining Long-Run Price Stability? – Conceptual Issues*, w: „Achieving Price Stability”, a symposium sponsored by the Federal Reserve Bank of Kansas City at Jackson Hole, Wyoming, August 29-31, 1996.
3. *Współczesna polityka pieniężna*, red. W. Przybylska-Kapuścińska, Difin, Warszawa 2008.
4. *Założenia polityki pieniężnej na 2004 r.* Narodowy Bank Polski, Warszawa, wrzesień 2003.

MONETARY POLICY RULES AND FORECASTING OF INFLATION RATE

Summary

The feedback exists between monetary policy rules and inflation rate forecast. In this paper we analyse the impact kind of monetary policy rules on inflation rate forecast. We apply Svensson model and the presentation in the form of state space to forecasting of inflation rate.