

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 56

3-4

2009



DOM WYDAWNICZY ELIPSA
WARSZAWA 2009

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA REDAKCYJNA

Andrzej S. Barczak, Czesław Domański, Krzysztof Jajuga, Witold Jurek,
Michał Kolupa (Przewodniczący), Józef Pociecha, Danuta Strahl

KOMITET REDAKCYJNY

Mirosław Szreder (Redaktor Naczelny),
Magdalena Osińska (Zastępca Redaktora Naczelnego),
Marek Walesiak (Zastępca Redaktora Naczelnego),
Krzysztof Najman (Sekretarz Naukowy)

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

Strona WWW „Przeglądu Statystycznego”:

<http://www.przegladstatystyczny.pan.pl>

Nakład 300 egz.



Realizacja wydawnicza:
Dom Wydawniczy ELIPSA,
ul. Inflancka 15/198, 00-189 Warszawa
tel./fax (0-22) 635 03 01, 635 17 85, e-mail: elipsa@elipsa.pl
www.elipsa.pl

Errata do nr 2/2009, tom 56

W artykule P. Kulczycki, K. Daniel, *Metoda wspomagania strategii marketingowej operatora telefonii komórkowej* powinno być:

1. s. 118 – w trzeciej i czwartej linii od dołu, „podziału zbioru danych – na przykład określonych w postaci próby losowej $x_1, x_2, \dots, x_m$ – na podgrupy” zamiast: „podziału zbioru danych – na przykład określonych w postaci próby losowej $x_1, x_2, \dots, x_m$ na podgrupy”,

2. s. 119 – w końcowej części wzoru (5) $k = 0, 1, \dots$, zamiast $k = 1, 1, \dots$,

3. s. 120 – w czwartej linii $m_1, m_2, \dots, m_J$, zamiast $m_1, m_2, \dots, m_j$, a także $\hat{f}_1, \hat{f}_2, \dots, \hat{f}_J$, zamiast $\hat{f}_1, \hat{f}_2, \dots, \hat{f}_j$,

4. na s. 133 – Instytut Badań **Systemowych** PAN, zamiast Instytut Badań Społecznych PAN, dwukrotnie w przypadku afiliacji obojga Autorów.

ANNA PAJOR, ARTUR PRĘDKI

ESTYMACJA MIERNIKA EFEKTYWNOŚCI TECHNICZNEJ W RAMACH METODY DEA¹

1. WSTĘP

Podstawowym zastosowaniem metody DEA jest pomiar efektywności technicznej jednostek gospodarczych, należących do grupy obiektów o zbliżonej technologii, która jest reprezentowana przez zbiór możliwości produkcyjnych (zob. [1], [2]). Możliwość ustalenia poziomu efektywności technicznej dla każdego obiektu z grupy pozwala utworzyć ranking tych jednostek oraz ustalić źródła ewentualnego braku ich efektywności. Głównymi zaletami metody DEA są (zob. [6], s. 58):

- brak konieczności parametrycznej specyfikacji zależności łączącej nakłady i produkty, czyli tzw. funkcji produkcji (dla jednego produktu) lub transformaty produkcji (dla wielu produktów, zob. [22]);
- możliwość jej zastosowania nawet przy dużej liczbie nakładów i produktów;
- prostota obliczeń (wykorzystanie programowania liniowego do obliczania wartości mierników);
- wartości nakładów i produktów mogą być wyrażone w różnych jednostkach, co więcej, zmiana jednostki danego czynnika (np. z kg na tonę) nie ma wpływu na wartość miernika (tzw. niezmienniczość miary efektywności).

Nie można jednak zapominać o jej zasadniczej wadzie, którą jest duża wrażliwość na błędne lub niekompletne dane (zob. [6], s. 59). Wynika to z deterministycznego charakteru metody, w której miernik efektywności obiektu liczony jest właśnie na podstawie danych dotyczących nakładów i produktów wybranej grupy obiektów. Co więcej, wybór owej grupy, jak i wybór nakładów i produktów użytych do analizy, jest często podyktowany dostępnością danych i zwykle nie ma pewności, że niosą one wystarczającą informację o opisywanej przez nas technologii. W deterministycznej wersji metody DEA postać zbioru możliwości produkcyjnych, reprezentującego technologię, również wyznacza się na podstawie dostępnych danych.

Zachodzi więc potrzeba modelowania owej niepewności oraz jej pomiaru (np. za pomocą miar rozproszenia, czy przedziałów ufności). Jednocześnie, chcąc zachować pierwszą z wymienionych zalet, chcielibyśmy uniknąć specyfikacji konkretnej parametrycznej zależności między nakładami a produktami. Celem niniejszej pracy jest prezentacja i ilustracja (pierwsza na gruncie polskim) propozycji modelu statystycznego

¹ Praca wykonana w ramach badań statutowych finansowanych przez Uniwersytet Ekonomiczny w Krakowie.

spełniającego owe wymagania². W części ósmej pracy przedstawiono własności rozważanych estymatorów dla małej próby w oparciu o symulację. Następnie zastosowano przedstawioną metodologię do analizy poziomu efektywności rzeczywistych obiektów produkcyjnych z polskiego sektora energetycznego.

2. METODA DEA

Rozważamy grupę n jednostek produkcyjnych, o których zakładamy, że:

- stosują tę samą technologię³,
- z p nakładów wytwarzają q produktów,
- dysponujemy danymi dotyczącymi ilości tych nakładów i produktów dla każdego obiektu w grupie.

Obiekty te można przedstawić w postaci (x_i, y_i) , $i = 1, \dots, n$, gdzie $x_i = [x_i^1, \dots, x_i^p] \in R_{+0}^p$ oraz $y_i = [y_i^1, \dots, y_i^q] \in R_{+0}^q$ oznaczają, odpowiednio, wielkość nakładów oraz produktów i -tego obiektu⁴. Niech $X_n = \{(x_i, y_i), i = 1, \dots, n\}$ oznacza zbiór tych obiektów. Wspólna technologia reprezentowana jest za pomocą zbioru możliwości produkcyjnych:

$$T = \{(x, y) \in R_{+0}^{p+q}: \text{z } x \text{ można wyprodukować } y\}. \quad (2.1)$$

Dla dowolnego, ustalonego $x \in R_{+0}^p$ zdefiniujemy też pomocniczo zbiór:

$$T(x) = \{y \in R_{+0}^q: (x, y) \in T\}. \quad (2.2)$$

W deterministycznej wersji metody DEA zakładamy, iż zbiór możliwości produkcyjnych T można przedstawić, za pomocą dostępnych danych, w postaci⁵:

$$T(X_n) = \left\{ (x, y) \in R_{+0}^{p+q}: y \leq \sum_{i=1}^n \lambda_i \cdot y_i, x \geq \sum_{i=1}^n \lambda_i \cdot x_i, \sum_{i=1}^n \lambda_i = 1, \forall i \in \{1, \dots, n\} \lambda_i \geq 0 \right\}. \quad (2.3)$$

Obiekt $(x_o, y_o) \in T$ nazywamy **efektywnym technicznie względem produktów**⁶, gdy niemożliwe jest (w sensie przynależności do zbioru T) proporcjonalne zwiększenie ilości jego produktów y_o i jednocześnie zachowanie ilości jego nakładów na dotychczasowym

² Autorami tej propozycji modelowania niepewności na gruncie metody DEA jest zespół naukowców skupiony wokół prof. Leopolda Simara (zob. prace [5], [8], [9], [12], [20]).

³ Założenie to zapewnia możliwość porównywania ze sobą tych jednostek np. pod względem poziomu efektywności technicznej. W zastosowaniach metody DEA staramy się więc, aby założenie to było, chociaż w przybliżeniu, spełnione. Wybieramy grupy obiektów z tej samej branży, sektora albo będące oddziałami większej jednostki gospodarczej (np. banku).

⁴ $R_{+0} = R_+ \cup \{0\}$.

⁵ Zauważmy, że zbiór ten jest wypukły i domknięty, zaś $T(x)$ jest wtedy ograniczony. Ponadto $\forall i \in \{1, \dots, n\} (x_i, y_i) \in T = T(X_n)$.

⁶ W dalszym ciągu pracy będziemy używać skróconej nazwy „obiekt efektywny technicznie”. Istnieje też tzw. efektywność techniczna względem nakładów (zob. np. [13]).

poziomie x_o . Miernikiem efektywności technicznej względem produktów dla obiektu (x_o, y_o) , gdzie $y_o \neq 0$, jest⁷:

$$\theta(x_o, y_o) = \sup \{ \theta \in R_+ : (x_o, \theta \cdot y_o) \in T \}, \quad (2.4)$$

zaś efektywny poziom produktów tego obiektu to:

$$y^\partial(x_o) = \theta(x_o, y_o) \cdot y_o. \quad (2.5)$$

Obiekt $(x, y^\partial(x))$ jest nazywany wzorcem efektywności obiektu (x, y) i nie musi być obserwowany w rzeczywistości.

W deterministycznej wersji metody DEA można też zapisać:

$$\theta(x_o, y_o) = \sup \{ \theta \in R_+ : (x_o, \theta \cdot y_o) \in T(\mathcal{X}_n) \}. \quad (2.6)$$

Miernik efektywności technicznej ma następujące własności:

- niezmienniczość względem przyjętych jednostek nakładów i produktów;
- $\forall (x_o, y_o) \in T \ \theta(x_o, y_o) \geq 1$;
- $(x_o, y_o) \in T$ jest efektywny technicznie $\Leftrightarrow \theta(x_o, y_o) = 1$.

Zwróćmy uwagę, że miernik ten jest liczony dla punktu ze zbioru T i zależy od tego zbioru. Może przyjmować różne wartości w zależności od wybranego obiektu (x_o, y_o) . Jest to więc lokalna charakterystyka technologii, która obliczana jest zwykle dla n obiektów z analizowanej grupy $\mathcal{X}_n$. Istota zastosowania metody DEA w tym przypadku polega na rozwiązaniu n zagadnień optymalizacyjnych z zakresu programowania liniowego:

$$\begin{aligned} P.1 \quad & \theta(x_i, y_i) \rightarrow \text{MAX} \\ & \theta(x_i, y_i) \cdot y \leq \sum_{j=1}^n \lambda_j \cdot y_j, \ x \geq \sum_{j=1}^n \lambda_j \cdot x_j, \ \sum_{j=1}^n \lambda_j = 1 \\ & \theta(x_i, y_i) \geq 0, \ \lambda_j \geq 0, \ \text{dla } j = 1, \dots, n. \end{aligned}$$

W wyniku uzyskujemy miernik efektywności dla każdego obiektu z grupy. Obiekty, dla których $\theta(x_i, y_i) > 1$ zostają uznane za nieefektywne, gdyż nie spełniają definicji obiektu efektywnego technicznie względem zbioru T zadanego wzorem 2.3. Sortując jednostki względem niemalejącej wartości tego miernika można otrzymać ich ranking od obiektów efektywnych do najbardziej nieefektywnych. Co więcej, interpretacja $\theta(x_i, y_i)$ dla jednostki nieefektywnej daje interesujące informacje o pożądanej strukturze produktów tej jednostki, która prowadziłaby do jej efektywności technicznej (zob. [21]). Zwróćmy jednak uwagę, że miernik ten ma charakter względny, tzn. efektywność danego obiektu mierzona jest względem wszystkich jednostek w grupie. Oznacza to w szczególności, że tak jak w przypadku zbioru możliwości produkcyjnych, wartość miernika zależy od danych, którymi dysponujemy.

⁷ Supremum istnieje ze względu na własności zbioru T (zob. przypis 5). Zapis $y_o \neq 0$ oznacza, że co najmniej jedna współrzędna wektora y_o jest niezerowa.

3. MODEL STATYSTYCZNY

Podsumowując niektóre uwagi poczynione w części drugiej można więc stwierdzić, iż zmiana zbioru X_n może prowadzić do zmiany postaci zbioru $T(X_n)$. Tym samym mogą zmienić się wartości mierników $\theta(x_i, y_i)$, liczone w oparciu o ten właśnie zbiór. Z sytuacją taką możemy mieć do czynienia, gdy np. część danych okaże się błędna lub nieprecyzyjna, albo gdy zdobędziemy dodatkowe dane poszerzające dotychczasowy zbiór X_n o kolejne obiekty. W ramach metody DEA próbowano rozważać ten problem na różne sposoby. Począwszy od tradycyjnej analizy wrażliwości w ramach programowania liniowego (zob. [3], [4]), poprzez różnego typu modele parametryczne i semiparametryczne (zob. [7], [14], [15]), aż do modeli nieparametrycznych (zob. [5], [9], [16], [20]). Zainteresowanie autorów niniejszej pracy wzbudziła szczególnie ta ostatnia kategoria ze względu na wyrażone we wstępie pragnienie zachowania dotychczasowych zalet metody DEA. Zalicza się do nich m.in. liczenie miernika efektywności głównie w oparciu o dostępne dane i unikanie stawiania, często arbitralnych, założeń o charakterze parametrycznym.

Wprowadzamy więc odpowiedni model statystyczny (za pracami [5], [9] i [20]). Definiują go kolejne założenia, których znaczenie będzie pokrótce komentowane.

Założenie 1

O zbiorze T zakładamy, iż:

- jest wypukły i domknięty,
- w przypadku braku nakładów nie można niczego wyprodukować⁸,
- spełnia warunek tzw. swobodnej dyslokacji nakładów i produktów⁹ tzn.:

$$\forall \underline{x} \geq x, \underline{y} \leq y [(x, y) \in T \Rightarrow (\underline{x}, \underline{y}) \in T], \quad (2.2)$$

– dla dowolnego, ustalonego $x \in R_{+0}^p$ zbiór $T(x) = \{y \in R_{+0}^q : (x, y) \in T\}$ jest ograniczony.

Założenie 2

Próbę $X_n = \{(X_i, Y_i), i = 1, \dots, n\}$ charakteryzuje ciąg niezależnych wektorów losowych $(X_i, Y_i): R_{+0}^{p+q} \rightarrow T$ o tym samym rozkładzie i gęstości (względem miary Lebesgue'a): $f_T(x, y): R_{+0}^{p+q} \rightarrow R$, spełniającej warunek: $\forall (x, y) \in T f_T(x, y) > 0$, zaś $\forall (x, y) \notin T f_T(x, y) = 0$, czyli $\text{supp}(f_T) = T$ (nośnik funkcji).

Zbiór X_n jest więc **próbą losową (prostą)** jednostek produkcyjnych. Natomiast zbiór możliwości produkcyjnych (T) jest nieznany¹⁰, ale może być estymowany. Zwróćmy

⁸ Bardziej formalnie oznacza to, że gdy x jest wektorem zerowym, natomiast y nie jest, to (x, y) nie należy do zbioru T .

⁹ Z ang. *free disposal assumption*. Nierówności między wektorami są tu rozumiane „po współrzędnych”.

¹⁰ Od tego momentu zbiór T nie może być utożsamiany ze zbiorem $T(X_n)$, jak to miało miejsce w wersji deterministycznej DEA. Wprowadzając zał. 1 zachowujemy jednak jego własności zawarte w przypisie 5.

uwagę, iż przestrzenią obserwacji, czyli zbiorem potencjalnych realizacji, jest tutaj R_{+0}^{p+q} . Nie możemy bowiem wykluczyć żadnej realizacji z tego nieujemnego ortantu, ze względu na to, iż postać zbioru T nie jest znana. Funkcja $f_T(x, y)$ jest zależna od postaci zbioru T i może być przedmiotem wnioskowania statystycznego.

Założenie 3¹¹

$f_T(x, y)$ jest ciągła na T oraz $\forall (x, y) \in \text{int}T: f(x, \theta(x, y) \cdot y) > 0$.

Warunek ten oznacza, iż istnieje niezerowe prawdopodobieństwo zaobserwowania obiektów leżących w bliskim otoczeniu wzorca efektywności.

Założenie 4

$\theta(x, y)$ jest funkcją klasy $C^2(T)$.

Założenie to gwarantuje odpowiednią „gładkość” funkcji $\theta(x, y)$.

Zdefiniowany model statystyczny będziemy oznaczać krótko przez $\mathcal{P}$ lub $\mathcal{P}(T, f_T(x, y))$. Niech $\mathfrak{T}$ oraz Ψ będą odpowiednio rodziną zbiorów T oraz gęstości f_T , spełniających założenia modelu. Zbiór $\mathfrak{T} \times \Psi$ jest wtedy odpowiednikiem przestrzeni parametrów w modelu parametrycznym. Podstawowym zadaniem wydaje się więc estymacja nieznanego zbioru T i funkcji f_T na podstawie próby $\mathcal{X}_n$. Z drugiej strony nie można zapominać, iż podstawowym zastosowaniem metody DEA jest możliwość obliczenia wartości miernika efektywności technicznej $\theta(x, y)$ dla ustalonego obiektu. Te dwa cele są jednak ze sobą ściśle związane. Miernik $\theta(x, y)$ (jako funkcja) występuje przecież w modelu statystycznym (zob. zał. 4).

4. ESTYMATOR MIERNIKA EFEKTYWNOŚCI TECHNICZNEJ

Próba realizacji celów postawionych pod koniec części trzeciej stanowi treść dalszych badań autorów w tym obszarze. W niniejszej pracy skupimy się na niektórych problemach związanych z oszacowaniem wartości funkcji $\theta(\cdot, \cdot)$ w dowolnym, ustalonym punkcie¹² $(x_o, y_o) \in R_{+0}^{p+q}$. Przypomnijmy w tym miejscu jej definicję (zob. wzór 2.4):

$$\theta(x_o, y_o) = \sup \{ \theta \in \mathbb{R}_+: (x_o, \theta \cdot y_o) \in T \}.$$

Zauważmy, że $\theta(x_o, y_o)$ zależy od nieznanego zbioru T .

Estymatorem nieznanego zbioru możliwości produkcyjnych T jest:

$$\hat{T}(\mathcal{X}_n) = \left\{ (x, y) \in R_{+0}^{p+q}: y \leq \sum_{i=1}^n \lambda_i \cdot Y_i, x \geq \sum_{i=1}^n \lambda_i \cdot X_i, \sum_{i=1}^n \lambda_i = 1, \forall i \in \{1, \dots, n\} \lambda_i \geq 0 \right\}. \quad (4.1)$$

¹¹ Symbol $\text{int}T$ oznacza wnętrze zbioru T .

¹² Cały czas obowiązuje założenie $y_o \neq 0$ (zob. przypis 7).

Warto przypomnieć, że w wersji deterministycznej DEA pojedyncza realizacja zbioru $\hat{T}(X_n)$ jest utożsamiana ze zbiorem $T(T(X_n))$, zob. wzór 2.3). W przypadku stochastycznej wersji tej metody tak nie jest (zob. przypis 10). Wykazano jednak, iż zbiory te są ze sobą ściśle związane. Po pierwsze, **realizacja** ciągu wektorów losowych (X_i, Y_i) (tj. (x_i, y_i) , $i = 1, \dots, n$) należy do zbioru T (zob. zał. 2). Realizacja zbioru $\hat{T}(X_n)$ jest niewątpliwie, na mocy swej konstrukcji, najmniejszym zbiorem spełniającym założenia modelu mówiące, że $(x_i, y_i) \in T(X_n)$ (dla $i = 1, \dots, n$) (zob. przypis 5). Skoro tak, to oczywiście realizacja $\hat{T}(X_n)$ zawiera się w zbiorze T . Po drugie, wykazano, czego nie będziemy tu szerzej omawiać, iż $\hat{T}(X_n)$ jest efektywnym estymatorem zbioru T w sensie ryzyka minimaxowego (szczegóły można znaleźć w pracach [10], [11]).

Estymatorem DEA¹³ wartości funkcji $\theta(\cdot, \cdot)$ w punkcie $(x_o, y_o) \in R_{+0}^{p+q}$ nazywamy statystykę $\hat{\theta}(x_o, y_o)$, zdefiniowaną następująco:

$$\hat{\theta}(x_o, y_o) = \sup\{\theta \in R_+ : (x_o, \theta \cdot y_o) \in \hat{T}(X_n)\}. \quad (4.2)$$

Skoro X_n jest ciągiem wektorów losowych, to $\hat{\theta}(x_o, y_o)$ jako ich funkcja (poprzez zbiór losowy $\hat{T}(X_n)$) może być nazywana statystyką i jako taka pełnić rolę estymatora wartości¹⁴ $\theta(x_o, y_o)$. Udowodniono zgodność tego estymatora (w każdym ustalonym punkcie $(x_o, y_o) \in R_{+0}^{p+q}$) oraz ustalono tzw. szybkość jego zbieżności według rozkładu (zob. [8]). W części piątej pracy zajmujemy się szczegółowo innym, ważnym, choć cząstkowym rezultatem. Mianowicie przy $p = q = 1$ (dla jednego nakładu i jednego produktu) wyprowadzono postać asymptotycznego rozkładu związanego ściśle z różnicą $\hat{\theta}(x_o, y_o) - \theta(x_o, y_o)$ (zob. [5]). Jeśli głównym zastosowaniem DEA jest ocena wartości miernika efektywności technicznej, to chcielibyśmy określić błąd związany z tym oszacowaniem (choćby asymptotyczny). Rezultat ten pozwala na konstrukcję asymptotycznych miar rozproszenia i odpowiednich przedziałów ufności.

5. ASYMPTOTYCZNY ROZKŁAD ESTYMATORA DEA

Na początek wprowadzimy pewne oznaczenia i notację. Przyjmujemy, w dalszym ciągu pracy, iż $p = q = 1$. Skoro zbiór T jest wypukły i $T(x)$ ograniczony (zob. zał. 1), to „ograniczenie górne” T może być opisane monotoniczną, mierzalną, wklęsłą funkcją $y = g(x)$ (funkcją produkcji, $g: R_{+0} \rightarrow R$), a więc:

$$T = \{(x, y) \in R_{+0}^2 : y \leq g(x)\}. \quad (5.1)$$

Z drugiej strony, dla danego $x_o \in R_{+0}$:

$$g(x_o) = \sup\{y \in R_{+0} : (x_o, y) \in T\}. \quad (5.2)$$

¹³ Z ang. *DEA-estimator*. W literaturze polskiej jednakże funkcjonuje w takiej sytuacji szyk przestawny (np. piszemy raczej estymator MNK, a nie MNK-estymator).

¹⁴ Skoro $T(X_n) \subseteq T$, to z definicji $\hat{\theta}(x_o, y_o) \leq \theta(x_o, y_o)$. Oczywiście, gdyby $T(X_n) = T$ to $\hat{\theta}(x_o, y_o) = \theta(x_o, y_o)$.

Estymatorem DEA $g(x_o)$ jest:

$$\hat{g}(x_o) = \sup\{y \in R_{+0} : (x_o, y) \in \hat{T}(\mathcal{X}_n)\}. \quad (5.3)$$

Dla $y_o \neq 0$ zachodzą proste, lecz istotne zależności:

$$\theta(x_o, y_o) = g(x_o)/y_o \text{ oraz } \hat{\theta}(x_o, y_o) = \hat{g}(x_o)/y_o. \quad (5.4)$$

Interesujący nas rezultat przedstawia następujące twierdzenie.

Twierdzenie 1

Niech $p = q = 1$ i spełnione są założenia modelu statystycznego wprowadzonego w części trzeciej pracy. Dodatkowo niech $g''(x_o) < 0$.

$$\text{Wówczas: } \forall z < 0: P\left\{n^{2/3} \cdot (b_0^2/b_2)^{1/3} \cdot [\hat{g}(x_o) - g(x_o)] \leq z\right\} = \int_0^\infty \varphi(u, z) du + o(1),$$

$$\text{gdzie } \varphi(u, z) = 0,5(-z)^{3/2} (1 + u^2) \exp[-1/6 \cdot (-z)^{3/2} \cdot (u + u^{-1})^3], \\ b_0 = f_T(x_o, g(x_o)), b_2 = -0,5g''(x_o) \text{ (nieznane parametry)}^{15}.$$

Zwróćmy uwagę, że druga pochodna funkcji g istnieje i jest niedodatnia (wynika to z założeń 1 i 4 modelu statystycznego oraz wzorów 5.4). Twierdzenie to przedstawiono w pracy [20], zaś dowód znajduje się w opracowaniu [5], gdzie powyższa postać jest modyfikacją twierdzenia głównego. Korzystając ze wspomnianych wzorów 5.4, można też zapisać tezę tego twierdzenia dla bardziej interesującej nas różnicy $\hat{\theta}(x_o, y_o) - \theta(x_o, y_o)$:

$$\forall z < 0: P\left\{n^{2/3} \cdot (b_0^2/b_2)^{1/3} \cdot y_o \cdot [\hat{\theta}(x_o, y_o) - \theta(x_o, y_o)] \leq z\right\} = \int_0^\infty \varphi(u, z) du + o(1). \quad (5.5)$$

Bardzo istotny jest również wniosek wynikający z twierdzenia 1, a mówiący o postaci asymptotycznych miar obciążenia i rozproszenia estymatora $\hat{g}(x_o)$.

Wniosek 1

Bas. $(\hat{g}(x_o)) = -n^{-2/3} (b_2/b_0^2)^{1/3} c_1$ (obciążenie),

Var^{as.} $(\hat{g}(x_o)) = n^{-4/3} (b_2/b_0^2)^{2/3} c_2$ (wariancja),

MSE^{as.} $(\hat{g}(x_o)) = n^{-4/3} (b_2/b_0^2)^{2/3} (c_1^2 + c_2)$ (błąd średniokwadratowy).

$$\text{gdzie } c_1 = \int_0^\infty \int_0^\infty \varphi(u, -z) du dz = 2 \cdot 6^{2/3} \cdot \Gamma(2/3)/9 \approx 0,99360$$

¹⁵ $o(1)$ to wyrażenie dążące do zera przy $n \rightarrow \infty$ (nie zależy od z). Brak kolejności w numeracji parametrów w tezie wynika z faktu, że w dowodzie twierdzenia występuje jeszcze jeden niejawny parametr $b_1 = g'(x_o)$ (zob. [5], s. 227).

$$c_2 = 2 \int_0^\infty \int_0^\infty z \varphi(u, -z) du dz - c_1^2 = 4 \cdot 6^{1/3} \cdot \Gamma(1/3)/15 - c_1^2 \approx 0,31088.$$

Znak obciążenia jest ujemny, ponieważ rzeczywista wartość funkcji produkcji $g(x_o)$ jest zawsze większa bądź równa od oceny¹⁶ $\hat{g}(x_o)$. Jeśli przeanalizujemy wyrażenie b_2/b_0^2 to widać, że na zwiększenie obciążenia lub wariancji może wpłynąć mniejsza wartość gęstości w punkcie $(x_o, g(x_o))$ (co można związać z mniejszą pewnością zaistnienia realnego obiektu w otoczeniu tego punktu) lub „większa” krzywizna g w punkcie x_o ($\hat{g}(x_o)$ jako funkcja x_o jest tylko przedziałami liniowa, więc nie jest zbyt „gładka”).

6. ESTYMACJA NIEZNANYCH PARAMETRÓW

Zarówno w twierdzeniu 1, jak i we wniosku 1 występują nieznanne parametry b_0 oraz b_2 . Zachodzi więc potrzeba ich estymacji, co wprowadza kolejny „szum” do asymptotycznych rezultatów przedstawionych w części piątej. Niech (x_i, y_i) , $i = 1, \dots, n$, oznacza realizację wektorów losowych (X_i, Y_i) , $i = 1, \dots, n$.

6.1. ESTYMACJA b_0

Dla dowolnego $x_o \in R_{+0}$ i $\delta > 0$ definiujemy zbiór:

$$S(\delta) = (x_o - \delta/2; x_o + \delta/2) \times R. \quad (6.1)$$

Następnie definiujemy zbiór:

$$D = S(\delta) \cap T(X_n) \cap \{(x, y): y \geq \hat{g}(x_o) - \delta/2\}. \quad (6.2)$$

Estymatorem $b_0 = f_T(x_o, g(x_o))$ jest:

$$\hat{b}_0 = \#\{(x_i, y_i) \in D\} \cdot [n\lambda(D)]^{-1}, \quad (6.3)$$

gdzie $\lambda(D)$ jest miarą Lebesgue’a zbioru D .

6.2 ESTYMACJA b_2

Dla dowolnego, ustalonego $h > 0$ oraz $w \in R$ niech:

l_w^L – oznacza odcinek łączący $(x_o, \hat{g}(x_o))$ z $(x_o - h/2, w)$,

l_w^R – oznacza odcinek łączący $(x_o, \hat{g}(x_o))$ z $(x_o + h/2, w)$.

¹⁶ Wynika to z definicji tych pojęć oraz wspomnianego już faktu, iż $T(X_n) \subseteq T$.

Definiujemy teraz odpowiednio:

$$\begin{aligned} Z^- &= \max \{w: \exists 1 \leq i \leq n: (x_i, y_i) \in l_w^L\} \text{ oraz } Z^+ = \\ &= \max \{w: \exists 1 \leq i \leq n: (x_i, y_i) \in l_w^R\}. \end{aligned} \quad (6.4)$$

A następnie:

$$Z_1^- = \min \{Z^-, \hat{g}(x_o - h/2)\} \text{ oraz } Z_1^+ = \min \{Z^+, \hat{g}(x_o + h/2)\}. \quad (6.5)$$

Gdy dla ustalonego h żaden punkt (x_i, y_i) nie leży w pasie $(x_o - h/2, x_o) \times R$, to Z^- nie jest określone. Wtedy przyjmujemy $Z_1^- = \hat{g}(x_o + h/2)$. Analogicznie, gdy w pasie $(x_o, x_o + h/2) \times R$ nie znajduje się żaden punkt (x_i, y_i) , to $Z_1^+ = \hat{g}(x_o - h/2)$.

W kolejnym kroku dobieramy wielomian drugiego stopnia g_1 przechodzący przez punkty:

$$(x_o - h/2, Z_1^-); (x_o, \hat{g}(x_o)); (x_o + h/2, Z_1^+).$$

Warto w tym miejscu podkreślić, że konstrukcja wielomianu drugiego stopnia, przechodzącego przez punkty:

$$(x_o - h/2, \hat{g}(x_o - h/2)); (x_o, \hat{g}(x_o)); (x_o + h/2, \hat{g}(x_o + h/2)),$$

nie jest skuteczna, gdyż w większości praktycznych przypadków punkty te są współliniowe.

Estymatorem $b_2 = -0,5g''(x_o)$ jest:

$$\hat{b}_2 = -0,5g_1''(x_o). \quad (6.6)$$

Zauważmy, że punkty $(x_o - h/2, Z_1^-), (x_o + h/2, Z_1^+)$ leżą „poniżej” wykresu funkcji $\hat{g}(\cdot)$ (poza przypadkiem współliniowości z $(x_o, \hat{g}(x_o))$, o którym wspomnieliśmy powyżej). Stąd, na mocy wklęsłości $\hat{g}(\cdot)$, wynika wklęsłość g_1 . Oznacza to, że znak estymatora $\hat{b}_2$ jest nieujemny.

Konstrukcjom opisanym w częściach 6.1 i 6.2 pracy nadaje sens kolejne twierdzenie.

Twierdzenie 2 (zob. [5], s. 223)

Niech spełnione będą założenia twierdzenia 1 oraz $\exists \varepsilon_1 \in (0, 1/3)$ $\exists \varepsilon_2 \in (0, 1/2)$:
 $h \sim n^{-1/3 + \varepsilon_1}$, $\delta \sim n^{-1/2 + \varepsilon_2}$.

Wówczas $\hat{b}_0, \hat{b}_2$ są zgodnymi estymatorami, odpowiednio, b_0 i b_2 .

7. ASYMPTOTYCZNY PRZEDZIAŁ UFNOŚCI

Na mocy przytoczonego wcześniej wniosku 1 skorygowany o obciążenie estymator $g(x_o)$ ma postać:

$$\tilde{g}(x_o) = \hat{g}(x_o) + n^{-2/3} (\hat{b}_2 / \hat{b}_0^2)^{1/3} c_1. \quad (7.1)$$

Asymptotyczny przedział ufności dla wartości $g(x_o)$ wynika z przybliżonej równości:

$$P[\hat{g}(x_o) - n^{-2/3} (\hat{b}_2 / \hat{b}_0^2)^{1/3} z_{1-\alpha/2} \leq g(x_o) \leq \hat{g}(x_o) + n^{-2/3} (\hat{b}_2 / \hat{b}_0^2)^{1/3} z_{\alpha/2}] \approx_{as.} 1 - \alpha \quad (7.2)$$

lub w wersji „skorygowanej”:

$$P[\tilde{g}(x_o) - n^{-2/3} (\hat{b}_2 / \hat{b}_0^2)^{1/3} (z_{1-\alpha/2} + c_1) \leq g(x_o) \leq \tilde{g}(x_o) + n^{-2/3} (\hat{b}_2 / \hat{b}_0^2)^{1/3} (z_{\alpha/2} + c_1)] \approx_{as.} 1 - \alpha. \quad (7.3)$$

Gdzie $z_{1-\alpha/2}$ oraz $z_{\alpha/2}$ oznaczają kwantyle odpowiedniego rozkładu z twierdzenia 1, tzn.:

$$\int_0^\infty \varphi(u, z_{\alpha/2}) du = \lim_{n \rightarrow \infty} P\{n^{2/3} \cdot (b_0^2 / b_2)^{1/3} \cdot [\hat{g}(x_o) - g(x_o)] < z_{\alpha/2}\} = \alpha/2, \quad (7.4)$$

$$\int_0^\infty \varphi(u, z_{1-\alpha/2}) du = \lim_{n \rightarrow \infty} P\{n^{2/3} \cdot (b_0^2 / b_2)^{1/3} \cdot [\hat{g}(x_o) - g(x_o)] \leq z_{1-\alpha/2}\} = 1 - \alpha/2. \quad (7.5)$$

Powyższe całki są przybliżane za pomocą metod całkowania numerycznego. Postać funkcji $\varphi(u, z)$ wyklucza analityczne metody obliczeń.

8. BADANIA SYMULACYJNE I EMPIRYCZNE

Celem prowadzonych symulacji jest analiza skończenie-próbkowych własności estymatorów $\tilde{g}(x_o)$, $\bar{g}(x_o)$ oraz ich zachowanie przy zwiększaniu liczebności próby. Wykorzystane zostaną, wprowadzone już wcześniej, miary rozproszenia, histogramy interesujących nas wielkości oraz stopień pokrycia nieznannej wartości $g(x_o)$ przez odpowiednie przedziały ufności. Na koniec przedstawione zostaną wyniki uzyskane dla danych rzeczywistych. Obliczenia wykonano przy użyciu autorskich procedur napisanych w programie Gauss 8.0.

8.1 WYNIKI SYMULACJI

Za pracą [5] rozważamy dwa modele symulacyjne.

MODEL 1

$$X \sim U[0,1], Y = g(X) \cdot \exp(-V), g(x) = x^{1/2}, V \sim \text{Exp}(3),$$

gdzie V i X są niezależnymi zmiennymi losowymi, natomiast $\text{Exp}(\lambda)$ oznacza rozkład wykładniczy ze średnią $1/\lambda$.

Przyjmujemy więc dla uproszczenia, iż nakład jedynego czynnika produkcji obiektu może przyjmować wartości z przedziału $[0,1]$, nie preferujemy jednak *a priori* żadnej z tych wartości. Funkcja g pełni rolę klasycznej, potęgowej funkcji produkcji. Przy czym nie tylko ona ma wpływ na obserwowaną wielkość produkcji obiektu. Innym, niezależnym czynnikiem jest tu nieefektywność techniczna, którą w modelu uwzględniono w postaci zmiennej losowej $\exp(-V)$. Jest on (dla każdej konkretnej realizacji zmiennych x_o, v_o) nie większy od jedności, a tym samym zrealizowana wielkość produkcji obiektu y_o nie jest większa niż wartość funkcji produkcji $g(x_o)$ (zob. np. [22] oraz wzór 5.1). Zmienna losowa V ma rozkład wykładniczy z parametrem odwrotności skali (λ) równym 3. Im wyższa wartość tego parametru, tym większe prawdopodobieństwo zaobserwowania wartości „w pobliżu” grafu funkcji produkcji (zob. komentarz do zał. 3). Fakt ten ma duże znaczenie dla dowodu zgodności estymatora $\hat{g}(x_o)$ oraz przy estymacji nieznanymi parametrów b_0, b_2 . Stąd dla porównania rozważamy również model, w którym wartość parametru odwrotności skali w rozkładzie wykładniczym jest dużo mniejsza.

MODEL 2

$$X \sim U[0,1], Y = g(X) \cdot \exp(-V), g(x) = x^{1/2}, V \sim \text{Exp}(1),$$

gdzie V i X są niezależnymi zmiennymi losowymi.

Dla każdego z modeli rozważamy próby 100, 500 oraz 1000 elementowe i na ich podstawie estymujemy wartość funkcji produkcji $g(\cdot)$ w trzech różnych punktach odcinka $[0,1]$, będącego dziedziną gęstości brzegowej zmiennej X ($x_o = 0,25, x_o = 0,5$ oraz $x_o = 0,75$). Wartości parametrów wygładzających przyjęto częściowo arbitralnie, lecz tak, by spełniały one założenia twierdzenia 2: $\delta = 2n^{-1/2}$, $h = 2n^{-1/3}$. Liczba wygenerowanych prób (symulacji) dla każdego z problemów estymacyjnych wynosi $N = 500$. We wszystkich przypadkach obliczono wartości średnie:

$$\bar{\hat{g}}(x_o) = \frac{1}{N} \sum_{s=1}^N \hat{g}^{(s)}(x_o) \text{ oraz } \bar{\tilde{g}}(x_o) = \frac{1}{N} \sum_{s=1}^N \tilde{g}^{(s)}(x_o), \quad (8.1)$$

gdzie $\hat{g}^{(s)}(x_o)$ i $\tilde{g}^{(s)}(x_o)$ oznaczają, odpowiednio, wartość estymatora DEA oraz jego wersji skorygowanej, obliczone w s -tej symulacji. Obliczono również błąd średniokwadratowy oraz odchylenie standardowe $\hat{g}(x_o)$ ze wzoru¹⁷:

$$\sqrt{\frac{1}{N(N-1)} \sum_{s=1}^N [\hat{g}^{(s)}(x_o) - \bar{\hat{g}}(x_o)]^2}. \quad (8.2)$$

¹⁷ Analogiczny wzór stosujemy dla wersji skorygowanej zastępując „daszek” „falką”.

Wyniki symulacji dla modeli 1 i 2 przedstawiono w tabeli 1. Zawiera ona w każdej komórce kolejno:

- średnie obciążenie estymatora, odpowiednio, $\hat{g}(x_o)$ i $\tilde{g}(x_o)$ (zob. wniosek 1),
- odchylenie standardowe estymatorów $\hat{g}(x_o)$ i $\tilde{g}(x_o)$,
- błąd średniokwadratowy estymatora, odpowiednio, $\hat{g}(x_o)$ i $\tilde{g}(x_o)$.

Tabela 1

Średnie obciążenie ($\times 10^{-2}$), odchylenie standardowe ($\times 10^{-4}$) i błąd średniokwadratowy ($\times 10^{-4}$) odpowiednich estymatorów

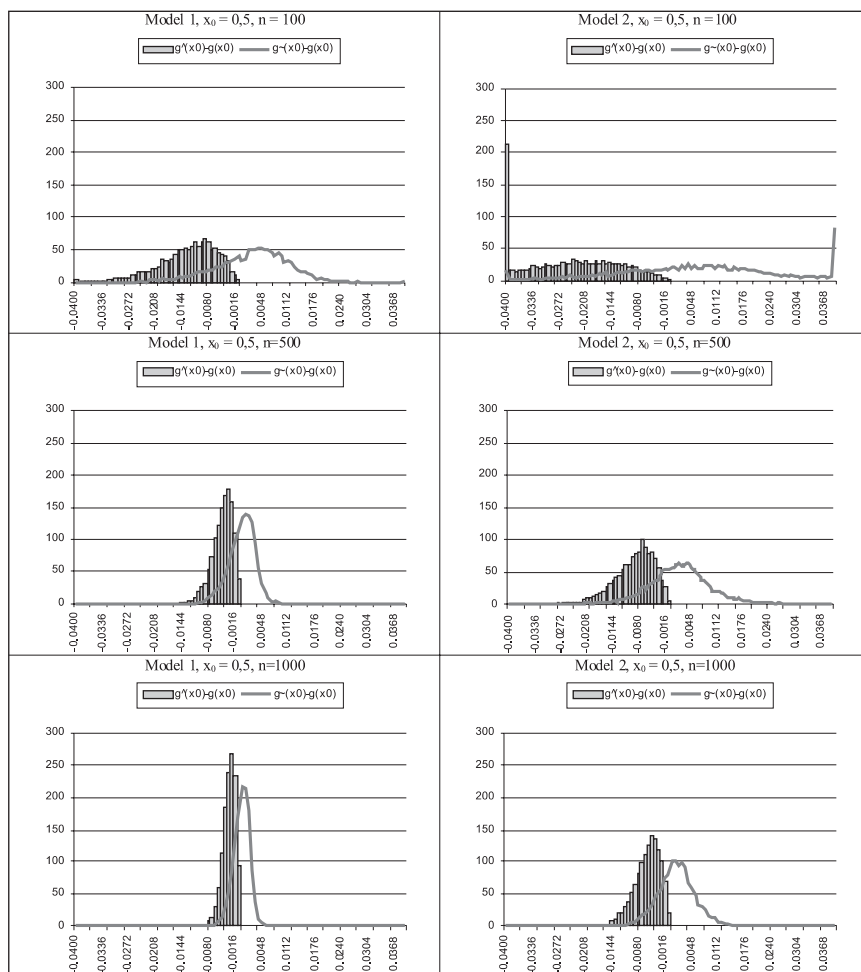
	$x_o = 0,25$				$x_o = 0,5$				$x_o = 0,75$			
	M1		M2		M1		M2		M1		M2	
	$\hat{g}$	$\tilde{g}$	$\hat{g}$	$\tilde{g}$	$\hat{g}$	$\tilde{g}$	$\hat{g}$	$\tilde{g}$	$\hat{g}$	$\tilde{g}$	$\hat{g}$	$\tilde{g}$
$n = 100$	-1,437	0,440	-2,960	0,235	-1,256	0,295	-2,638	0,560	-1,285	0,500	-3,101	0,863
	3,675	4,325	6,930	8,908	3,102	4,398	6,455	9,312	3,488	4,555	9,422	12,390
	2,738	1,127	11,157	4,015	2,058	1,052	9,041	4,640	2,260	1,286	14,046	8,405
$n = 500$	-0,449	0,080	-1,009	0,054	-0,422	0,064	-0,889	0,265	-0,398	0,081	-0,869	0,517
	1,131	1,367	2,479	3,468	1,034	1,361	2,118	3,205	1,006	1,496	2,317	3,772
	0,266	0,100	1,325	0,603	0,232	0,096	1,015	0,583	0,209	0,118	1,023	0,977
$n = 1000$	-0,299	0,020	-0,603	0,067	-0,268	0,040	-0,555	0,181	-0,261	0,055	-0,529	0,305
	0,778	0,867	1,433	1,905	0,649	0,857	1,437	2,136	0,645	0,921	1,393	2,298
	0,119	0,038	0,466	0,185	0,093	0,038	0,411	0,261	0,089	0,045	0,376	0,357

Źródło: opracowanie własne.

Po pierwsze, zwróćmy uwagę, że wraz ze wzrostem liczby obserwacji (n) odpowiednie miary rozproszenia oraz obciążenie zbliżają się do 0. Potwierdza to zgodność estymatorów $\hat{g}(x_o)$ i $\tilde{g}(x_o)$ (zob. tw. 1). Po drugie, widoczne jest, iż zarówno średnie obciążenie, rozproszenie, jak i błąd średniokwadratowy rozważanych estymatorów w modelu 2 są, co do wartości bezwzględnej, większe niż w modelu 1. Potwierdza to wzmiankowany wcześniej fakt, że estymacja w modelu 2 jest mniej precyzyjna ze względu na dużo mniejszą liczbę obserwacji położoną w pobliżu wykresu funkcji produkcji g . Po trzecie, zwróćmy uwagę, że korekta estymatora $\hat{g}(x_o)$ o obciążenie powoduje zmniejszenie średniego obciążenia i błędu średniokwadratowego. Zwiększa się jednak odchylenie standardowe $\tilde{g}(x_o)$. Oznacza to, że korekta o obciążenie powoduje zwiększenie rozproszenia ocen $g(x_o)$. To zwiększone rozproszenie wynika z konieczności szacowania nieznanymi parametrów b_0 , b_2 . Po czwarte, widoczna jest zależność wartości odpowiednich estymatorów od krzywizny funkcji produkcji. Dla funkcji produkcji postaci $g(x) = x^{1/2}$ krzywizna jest największa w pobliżu zera. Wraz ze wzrostem krzywizny¹⁸, zwiększeniu (co do wartości bezwzględnej) ulegają wszystkie trzy wielkości zilustrowane w tabeli 1. Wystarczy porównać te wielkości dla $x_o = 0,25$ (stosunkowo

¹⁸ Mierzonej wartością drugiej pochodnej g w odpowiednim punkcie.

blisko 0) z odpowiadającymi im wartościami dla $x_o = 0,5$ lub $0,75$. Zwróćmy jednak uwagę, że gdybyśmy chcieli rozróżnić, pod tym kątem, wyniki dla $x_o = 0,5$ oraz $x_o = 0,75$ sprawa już nie jest tak klarowna. Dopiero przy większej liczbie obserwacji (u nas równej 500 lub 1000) uwidoczni się rola krzywizny¹⁹. Wynika to z ogólnie znanego faktu, iż wraz ze wzrostem x_o krzywizna funkcji $x^{1/2}$ maleje, początkowo gwałtownie (blisko zera), lecz potem spadek ten nie jest już tak gwałtowny. Po piąte, estymator $\hat{g}(x_o)$ oczywiście niedoszacowuje rzeczywistej wartości $g(x_o)$ (wartość obciążenia jest ujemna). Sytuacja zmienia się po korekcie o obciążenie. Następuje przeszacowanie wartości $g(x_o)$, ale dużo mniejsze (co do wartości bezwzględnej) niż wcześniejsze niedoszacowanie.



Rysunek 1. Histogramy rozkładów $\hat{g}(0,5) - g(0,5)$ oraz $\tilde{g}(0,5) - g(0,5)$

Źródło: opracowanie własne.

¹⁹ Charakterystyki estymatorów dla $x_o = 0,5$ są większe, co do wartości bezwzględnej, od tych samych charakterystyk dla $x_o = 0,75$, ale tylko dla $n = 500$ lub 1000 .

W celu pełniejszej ilustracji wyników, omówionych w poprzednim akapicie, przedstawiamy histogramy rozkładów $\hat{g}(0,5) - g(0,5)$ oraz $\tilde{g}(0,5) - g(0,5)$. Ich postać jest widoczna na rysunku 1, w zależności od liczby obserwacji (n) i rodzaju modelu. W tym przypadku przeprowadzono 5000 symulacji.

Zwróćmy uwagę, że wraz ze wzrostem liczby obserwacji w obu modelach histogramy stają się bardziej wysmukłe i mniej rozproszone. Potwierdza to uwagi poczynione uprzednio, a związane ze zgodnością odpowiednich estymatorów. Ponadto warto też zauważyć, iż w modelu 2 odpowiednie histogramy są bardziej rozproszone i mniej skupione wokół modalnej niż w modelu 1. Potwierdza to uwagi o trudnościach w estymacji w ramach modelu 2. Z rysunków widać, iż histogramy rozkładu estymatora skorygowanego $\tilde{g}(x_o)$ są bardziej rozproszone niż odpowiednie histogramy rozkładu estymatora $\hat{g}(x_o)$. We wszystkich przypadkach, z powodu korekty o obciążenie, są też bardziej przesunięte w prawo. Korekta ta powoduje, jak już było wspomniane, większe rozproszenie estymatora $\tilde{g}(x_o)$.

Obliczono również realizacje asymptotycznych przedziałów ufności dla $g(x_o)$, przy poziomie istotności $\alpha = 0,05$ i $x_o = 0,5$ (zob. wzór 7.3). Rozważono oba modele i wykonano 500, a następnie 5000 symulacji, równolegle biorąc pod uwagę próby 100, 500 oraz 1000 elementowe. Pokrycie rzeczywistej wartości $g(0,5)$ przedstawiono w tabeli 2.

Tabela 2

Pokrycie rzeczywistej wartości $g(0,5)$ przez realizacje asymptotycznego przedziału ufności dla $\alpha = 0,05$

	$N = 500$		$N = 5000$	
	M1	M2	M1	M2
$n = 100$	0,906	0,902	0,913	0,896
$n = 500$	0,908	0,906	0,891	0,895
$n = 1000$	0,910	0,910	0,898	0,895

Źródło: opracowanie własne.

Nie zaobserwowano wzrostu stopnia pokrycia zarówno wraz ze wzrostem liczby obserwacji, jak i wraz ze wzrostem liczby symulacji. Szczególnie niepokojące jest to z punktu widzenia liczby obserwacji. Skoro są to realizacje **asymptotycznych** przedziałów ufności, być może zbieżność pokrycia do teoretycznej wartości 0,95 następuje przy dużo większej liczbie obserwacji²⁰. Powodem może być też niepewność wynikająca z szacowania nieznanymi parametrów b_0 , b_2 . Nie zaobserwowano również, aby stopień pokrycia $g(x_o)$ w modelu 1, teoretycznie łatwiej estymowalnego, był większy niż w modelu 2. Przeprowadzono podobne symulacje dla $x_o = 0,25$ oraz $x_o = 0,75$, stopień pokrycia $g(x_o)$ był zbliżony. Wydaje się więc, że krzywizna funkcji produkcji g nie odgrywa w tym przypadku znaczącej roli.

²⁰ Liczba obserwacji większa niż 1000 jest i tak nieużyteczna z punktu widzenia praktycznego zastosowania metody DEA.

8.2 ILUSTRACJA EMPIRYCZNA

Jako ilustrację empiryczną przedstawiamy wyniki uzyskane dla danych rzeczywistych z roku 1995, dotyczących 32 polskich elektrowni i elektrociepłowni. Jest to grupa jednostek produkcyjnych, dla których przeanalizowano stopień efektywności technicznej i jego rozproszenie. Jako jedyny nakład przyjęto kapitał liczony wartością brutto środków trwałych (x_i w zł)²¹. Produktem działalności jednostek jest natomiast wytworzona energia (y_i liczona w TJ²²). Dla uproszczenia analizy do obliczeń wzięto logarytmy tych wielkości²³ i przyjęto oznaczenia $\underline{x}_i = \ln x_i$ oraz $\underline{y}_i = \ln y_i$ ($i = 1, \dots, 32$). Przyjmując za x_o kolejne $\underline{x}_i$ (jak to najczęściej bywa w praktycznych zastosowaniach DEA) i korzystając z programu liniowego P.1, obliczono wartości estymatora miary efektywności technicznej dla poszczególnych obiektów. Następnie, w oparciu o zależności opisane we wzorach 5.4 i 7.1, obliczono wartości estymatorów $\hat{g}(x_o)$ i $\bar{g}(x_o)$. Wyniki te zebrano w tabeli 3, gdzie umieszczono również logarytmy rozważanych danych oraz realizacje 95% asymptotycznych przedziałów ufności (zob. wzór 7.3).

Tabela 3

Logarytm wielkości kapitału ($\underline{x}_i$) i produkcji ($\underline{y}_i$), wartości odpowiednich estymatorów, końce asymptotycznych przedziałów ufności

i	$\underline{x}_i$	$\underline{y}_i$	$\hat{\theta}(\underline{x}_i, \underline{y}_i)$	$\hat{g}(\underline{x}_i)$	$\bar{g}(\underline{x}_i)$	L	P
1	2	3	4	5	6	7	8
1	10,99705	8,056141	1	8,056141	x	x	x
2	11,08609	7,569928	1,077355	8,155497	x	x	x
3	11,26047	8,107991	1,029858	8,350079	8,867839	8,440072	9,546514
4	11,45214	8,275529	1,034853	8,563958	9,034749	8,645787	9,651858
5	11,62027	8,394958	1,042479	8,751563	9,25235	8,838606	9,908776
6	11,83614	8,497133	1,058292	8,992446	9,28595	9,043461	9,670672
7	12,00512	8,405703	1,092235	9,181003	9,45642	9,228873	9,817434
8	12,16236	8,419338	1,111306	9,356459	9,701397	9,416413	10,15354
9	12,23909	9,22219	1,023844	9,44208	9,834953	9,510367	10,34993
10	12,32511	9,453514	1,008944	9,538068	9,969169	9,612999	10,53425
11	12,6791	8,784162	1,130793	9,933069	10,02393	9,948862	10,14302
12	12,68296	8,977412	1,106931	9,937379	10,03028	9,953526	10,15205

²¹ Dane wielonakładowe (dostępne np. w pracy [13]). Ze względu na założenie: $p = q = 1$, rozważania ograniczono do jednego nakładu.

²² 1GWh = 3,6TJ (teradżul).

²³ Analogicznie zresztą jak w źródłowej pracy [5]. Zmniejsza to znacząco rozproszenie wartości danych i pozwala lepiej zilustrować działanie metody DEA w ramach wprowadzonego modelu statystycznego.

cd. tabeli 3

i	x_i	y_i	$\hat{\theta}(x_i, y_i)$	$\hat{g}(x_i)$	$\bar{g}(x_i)$	L	P
1	2	3	4	5	6	7	8
13	12,71478	9,228681	1,08064	9,972882	10,06739	9,989309	10,19127
14	12,84077	9,591588	1,113835	10,13698	10,36912	10,17733	10,67341
15	12,84363	9,63683	1,049791	10,11666	10,34765	10,15681	10,65042
16	12,86184	9,100972	1,054411	10,11347	10,34437	10,1536	10,64703
17	12,94949	9,149156	1,118658	10,23478	10,51175	10,28292	10,8748
18	12,99786	9,706785	1,059956	10,28876	10,56579	10,33691	10,92892
19	13,21364	9,911476	1,062359	10,52954	10,73523	10,56529	11,00485
20	13,28765	9,718404	1,091961	10,61212	10,83261	10,65044	11,12162
21	13,55388	10,9092	1,116817	10,93488	11,32771	11,00316	11,84262
22	13,62764	9,791114	1	10,9092	11,3192	10,98046	11,85662
23	13,76503	10,44455	1,067569	10,98411	11,28095	11,0357	11,67006
24	13,769	10,2889	1,051526	10,98272	11,27896	11,03421	11,66727
25	13,93473	10,37862	1,063901	11,04181	11,38241	11,10101	11,82885
26	13,95774	10,42618	1,059815	11,04983	11,41651	11,11356	11,89714
27	14,05469	10,25895	1,080383	11,08359	11,46566	11,15	11,96649
28	14,15775	10,84933	1,0249	11,11947	11,53554	11,19179	12,08092
29	14,31546	10,26171	1,08894	11,17439	11,38152	11,21039	11,65303
30	14,67309	10,43387	1,082908	11,29892	11,655	11,36081	12,12175
31	15,3091	11,52039	1	11,52039	12,47837	11,6869	13,73409
32	15,71262	9,889977	1,164855	11,52039	11,94631	11,59442	12,5046

Źródło: opracowanie własne.

Obiekty o nr 1, 22 i 31 zostały uznane za efektywne technicznie, ponieważ $\hat{\theta}(x_i, y_i) = 1$.

Dla pozostałych obiektów miara ta jest również stosunkowo bliska jedności, co oznacza, że są one wprawdzie nieefektywne, ale w niewielkim stopniu. Przykładowo, dla jednostki nr 32 wartość miary $\hat{\theta}(x_{32}, y_{32})$ jest największa i wynosi 1,164855. Oznacza to, że logarytm wielkości produkcji można zwiększyć o ponad 16% bez zmiany obecnej wartości kapitału. Największe, technologicznie możliwe do uzyskania, logarytmy produkcji dla poszczególnych obiektów to $\hat{g}(x_i)$, czyli kolumna piąta tabeli 3. Na przykład dla obiektu nr 32 $\hat{g}(x_{32}) = 11,52039 = y_{32} \cdot \hat{\theta}(x_{32}, y_{32})$. Powstaje więc ona poprzez wymnożenie kolumny 3 i 4 po wierszach (zgodnie ze wzorami 5.4).

W kolejnych kolumnach tabeli przedstawiono wyniki związane z zastosowaniem estymatora DEA. Kolumna szósta tabeli to „optymalne” wartości logarytmów produkcji z kolumny piątej skorygowane o ocenę obciążenia liczoną według wzoru z wniosku 1. Obciążenie to ma znak ujemny, więc wartości $\tilde{g}(\underline{x}_i)$ są przesunięte „na prawo” w stosunku do wartości $\hat{g}(\underline{x}_i)$. Przy obliczaniu oceny obciążenia popełniamy błędy związane z szacowaniem nieznanymi parametrów b_0 i b_2 . W tym miejscu należy zwrócić uwagę na stosunkowo arbitralny wybór stałych δ i h . W twierdzeniu 2 zawarte są zalecenia co do ich wyboru, zapewniające zgodność estymatorów $\hat{b}_0$ i $\hat{b}_2$, ale dla tak małej próby nie da się ich zrealizować. Przy wielkościach δ i h , sugerowanych we wspomnianym twierdzeniu²⁴, może się bowiem okazać, że dla niektórych obiektów z próby zbiór D jest pusty (zob. wzór 6.2). Jest to spowodowane stosunkowo dużą „rzadkością” chmury danych empirycznych²⁵. Niestety, niepustość tego zbioru jest warunkiem koniecznym dla obliczenia wartości estymatora $\hat{b}_0$. Stąd zarówno w źródłowej pracy [5], jak i w tym opracowaniu stałe te dobrano tak, by były one jak najbliżej przedziału (0,1) (zob. przypis 24), a jednocześnie tak, by obliczenia były wykonalne. Co więcej, obiekty nie są rozłożone równomiernie (zob. rys. 2), stąd konieczność nie tylko arbitralnego doboru stałych, ale także zmiany ich wartości w zależności od tego, w którym „regionie” danych się znajdujemy. Ostatecznie przyjęto²⁶:

$$\begin{aligned}\delta &= 2, h = 0,5 \text{ dla } i = 3,4, \\ \delta &= 2, h = 1 \text{ dla } i = 5, \dots, 25, \\ \delta &= 4, h = 2 \text{ dla } i = 26, \dots, 32.\end{aligned}$$

Niestety, nie udało się znaleźć sensownych (bliskich 1) wartości δ i h dla obiektów²⁷ o nr 1 i 2, stąd brak niektórych wartości w tabeli 3.

W ostatnich dwóch kolumnach tabeli 3 znajdują się wartości lewych i prawych końców realizacji asymptotycznych 95% przedziałów ufności dla rzeczywistych wartości $g(\underline{x}_i)$. Dają one pewną informację, ile mogłaby wynosić dla danego obiektu optymalna wartość funkcji produkcji – szczególnie wobec faktu, że przedziały te okazały się stosunkowo wąskie. Maksymalna różnica $P - L$ dla obiektu nr 31 wynosi około 2,047, co stanowi nieco ponad 17,5% wartości skorygowanego estymatora²⁸ $\tilde{g}(\underline{x}_{31})$. Do obliczonych realizacji należy jednak podchodzić z dużą ostrożnością. Po pierwsze, przy tak małej próbie korzystanie ze wzorów o charakterze asymptotycznym może skutkować dość dużym błędem. Po drugie, wspomniany już arbitralny dobór stałych δ i h wpływa na ocenę obciążenia, a tym samym na szerokość odpowiednich przedziałów.

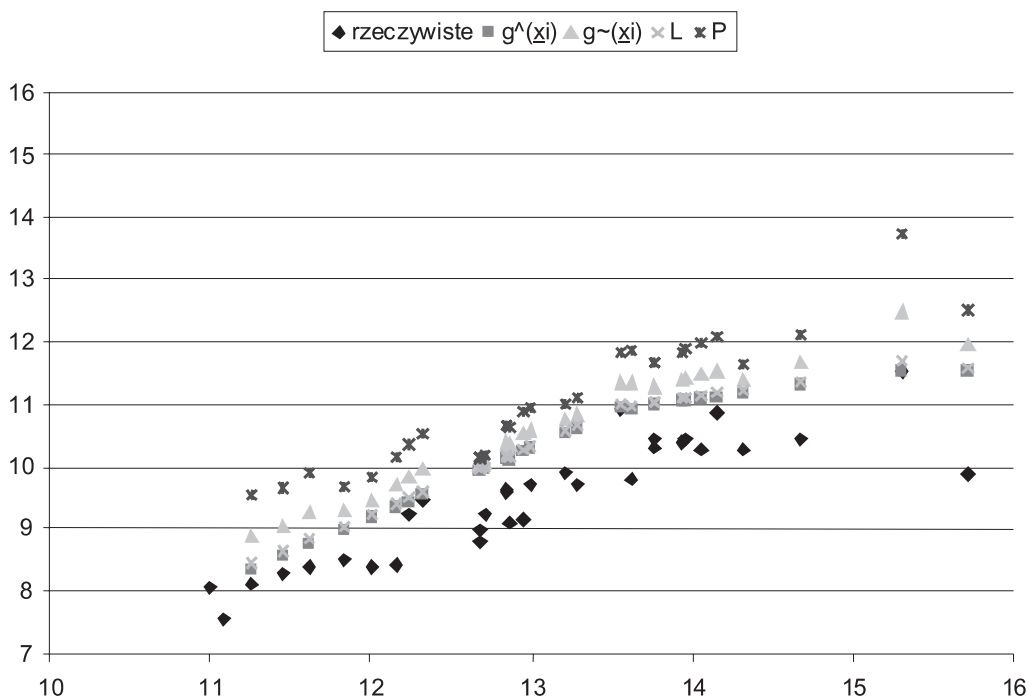
²⁴ Z tw. 2 wynika m.in., iż δ i h powinny należeć do przedziału (0,1).

²⁵ Stosunkowo duże odległości między obiektami powodują, że pewne zbiory pełniące rolę topologicznych otoczeń są puste.

²⁶ Dane w tabeli ponumerowane są wg rosnącego logarytmu kapitału. Regiony są więc tworzone właśnie względem tej wartości.

²⁷ Ten sam problem występuje w źródłowej pracy [5].

²⁸ Jest to oczywiście mocno nieformalny pomiar precyzji przedziałów ufności.



Rysunek 2. Wyniki zastosowania estymatora DEA w badaniu efektywności technicznej jednostek produkcyjnych

Źródło: opracowanie własne.

Patrząc na rysunek 2 widzimy, że realizacje przedziałów ufności są nieco węższe w środkowym „regionie”, gdzie punkty empiryczne są bliżej siebie (odpowiada to mniej więcej regionowi dla $i = 5, \dots, 25$, gdzie $\delta = 2$, $h = 1$). Zwróćmy też uwagę, że wartości estymatorów $\hat{g}(x_i)$ leżą poza skorygowanymi realizacjami przedziałów ufności. Łatwo przeliczyć, że ta sama sytuacja występuje dla nieskorygowanych realizacji (zob. wzór 7.2). Jest to konsekwencja ujemności kwantyli $z_{0,975}$, $z_{0,025}$. W głębszym zaś sensie ponownie wracamy do uwag poczynionych pod wnioskiem 1, mówiących o tym, że rzeczywista wartość $g(x_i)$ nie może leżeć poniżej wartości $\hat{g}(x_i)$, zaś rozbieżność między tymi wartościami świadczy o występowaniu nieefektywności technicznej.

9. PODSUMOWANIE

W niniejszej pracy przedstawiono propozycję modelowania niepewności dotyczącej postaci zbioru możliwości produkcyjnych oraz szacowania stopnia nieefektywności technicznej w ramach metody DEA, z zachowaniem jednocześnie zalet tej metody (o których była mowa we wstępie pracy). Propozycja ta jest nieznaną szerzej na gruncie polskim, ale obecna w literaturze światowej (zob. [5], [8], [9], [20]). Z przedstawionych w tej

pracy wstępnych badań wynika, że ze względu na asymptotyczny charakter własności stosowanych estymatorów prezentowana metoda daje dobre wyniki w przypadku dużej liczby obserwacji (zob. część symulacyjna artykułu), choć i wtedy nie jest pozbawiona usterek (np. niepełny stopień pokrycia przedziałów ufności, zob. tabela 2). Niestety, w praktycznych zastosowaniach metody DEA mamy najczęściej do czynienia ze zbiorami danych, które nie przekraczają 100 czy 200 obserwacji. W takiej sytuacji pojawia się kolejny problem związany z obliczaniem oceny obciążenia, a dokładniej, dotyczący wykonalności estymacji nieznanymi parametrów b_0 i b_2 oraz zbyt arbitralnego i dyskusyjnego wyboru pewnych stałych służących wspomnianej estymacji.

W małych próbach, przy braku jakichkolwiek rezultatów teoretycznych, postuluje się więc użycie metod bootstrapowych mających za zadanie „naśladować” rzeczywisty proces generujący dane. W literaturze przedmiotu nie brakuje propozycji w tym zakresie (zob. prace [17], [19], [20]). Niestety, do roku 2007 nie udowodniono nawet zgodności tych procedur²⁹, udowodniono natomiast niezgodność niektórych procedur zaproponowanych w latach poprzednich (zob. np. [18]). W roku 2008 pojawiła się jednak praca [9], w której zaproponowano pewne zgodne procedury bootstrapowe. Co więcej, przedstawiono w niej postać asymptotycznego rozkładu dla estymatora DEA w sytuacji wielu nakładów i produktów. Jej przeanalizowanie i zastosowanie będzie następnym zadaniem stojącym przed autorami niniejszego opracowania.

Kolejnym wyzwaniem będzie uzyskanie podobnych wniosków (wyników) dla estymatora miernika efektywności względem nakładów (zob. np. [20]). W niniejszej pracy badamy bowiem tzw. efektywność względem produktów (zob. wzór 2.4). W celach porównawczych i dla osiągnięcia szerszej perspektywy spojrzenia na problem modelowania niepewności w ramach metody DEA, należałoby również zbadać i wykorzystać inne propozycje modeli o charakterze semiparametrycznym i parametrycznym (zob. prace [7], [14], [15]).

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] Banker R.D., Charnes A., Cooper W.W., [1984], *Some models for estimating technical and scale inefficiencies in DEA*, Management Science, Vol. 30 (No. 9), s. 1078-1091.
- [2] Charnes A., Cooper W.W., Rhodes E., [1978], *Measuring the efficiency of decision making units*, European Journal of Operations Research, Vol. 2, s. 429-444.
- [3] Charnes A., Neralić L., [1989], *Sensitivity analysis in Data Envelopment Analysis – Part 1 i 2*, Glasnik Matematički, Vol. 24 (No. 44), s. 211-226 oraz 449-463.
- [4] Charnes A., Neralić L., [1992], *Sensitivity analysis in Data Envelopment Analysis – Part 3*, Glasnik Matematički, Vol. 27 (No. 47), s. 191-201.
- [5] Gijbels I., Mammen E., Park B.U., Simar L., [1999], *On estimation of monotone and concave frontier function*, „Journal of the American Statistical Association”, Vol. 94, s. 220-228.
- [6] Gospodarowicz M., [2000], *Procedury analizy i oceny banków*, Materiały i Studia NBP.
- [7] Grosskopf S., [1996], *Statistical inference and nonparametric efficiency: a selective survey*, „Journal of Productivity Analysis”, Vol. 7, s. 161-176.

²⁹ Własności podstawowej procedury bootstrapowej oznaczającej, że naśladuje ona w miarę dokładnie generowanie danych przez prawdziwy proces, przynajmniej asymptotycznie (przy dużej liczbie obserwacji).

- [8] Kneip A., Park B.U., Simar L., [1998], *A Note on the convergence of nonparametric DEA estimators for production efficiency scores*, *Econometric Theory*, Vol. 14, s. 783-793.
- [9] Kneip A., Simar L., Wilson P.W., [2008], *Asymptotics and consistent bootstraps for DEA estimators in nonparametric Frontier Models*, *Econometric Theory*, Vol. 24, s. 1663-1697.
- [10] Korostelev A.P., Simar L., Tsybakov A.B., [1995], *Efficient estimation of monotone boundaries*, „*The Annals of Statistics*”, Vol. 23, s. 476-489.
- [11] Korostelev A.P., Simar L., Tsybakov A.B., [1995], *On estimation of monotone and convex boundaries*, *Publications de l'Institut de Statistique des Universités de Paris XXXIX*, Vol. 1, s. 3-18.
- [12] Park B.U., Simar L., Weiner Ch., [2000], *The FDH estimator for productivity efficiency scores*, *Econometric Theory*, Vol. 16, s. 855-877.
- [13] Prędko A., [2003], *Analiza efektywności za pomocą metody DEA: podstawy formalne i ilustracja ekonomiczna*, „*Przegląd Statystyczny*”, Vol. 50 (No. 1), s. 87-100.
- [14] Sengupta J.K., [1990], *Transformations in stochastic DEA models*, „*Journal of Econometrics*”, Vol. 46, s. 109-123.
- [15] Simar L., [1992], *Estimating efficiencies from Frontier Models with panel data: a comparison of parametric, non-parametric and semi-parametric methods with bootstrapping*, „*Journal of Productivity Analysis*”, Vol. 3, s. 167-203.
- [16] Simar L., [1996], *Aspects of statistical analysis in DEA-type Frontier Models*, „*Journal of Productivity Analysis*”, Vol. 7, s. 177-185.
- [17] Simar L., Wilson P.W., [1998], *Sensitivity analysis of efficiency scores: how to bootstrap in nonparametric frontier models*, *Management Science*, Vol. 44, s. 49-61.
- [18] Simar L., Wilson P.W., [1999], *Some Problems with the Ferrier/Hirschberg Bootstrap Idea*, „*Journal of Productivity Analysis*”, Vol. 11, s. 67-80.
- [19] Simar L., Wilson P.W., [2000], *A general methodology for bootstrapping In nonparametric frontier models*, „*Journal of Applied Statistics*”, Vol. 27, s. 779-802.
- [20] Simar L., Wilson P.W., [2000] *Statistical Inference in Nonparametric Frontier Models: The State of the Art*, „*Journal of Productivity Analysis*”, Vol. 13, s. 49-78.
- [21] Sueyoshi T., [1997], *Measuring Efficiencies and Returns to Scale of Nippon Telegraph...*, *Management Science*, Vol. 43 (No. 6), s. 779-796.
- [22] Varian H.R., [1992], *Microeconomic Analysis*, W.W. Norton & Company, Inc.

Praca wpłynęła do redakcji w lipcu 2009 r.

ESTYMACJA MIERNIKA EFEKTYWNOŚCI TECHNICZNEJ W RAMACH METODY DEA

Streszczenie

W artykule przedstawiono propozycję opisu niepewności, co do wartości miernika efektywności technicznej, uzyskiwanej za pomocą metody DEA. Wprowadzono nieparametryczny model statystyczny, funkcjonujący w ramach omawianej metody oraz zdefiniowano tzw. estymator DEA miernika efektywności technicznej. W przypadku jednego nakładu i jednego produktu powiązano wartość miernika efektywności technicznej z wartością funkcji produkcji w punkcie oraz zdefiniowano odpowiadający jej estymator DEA. Następnie przedstawiono postać asymptotycznego rozkładu obu estymatorów DEA. Znajomość tego rozkładu pozwoliła wyznaczyć asymptotyczne obciążenie i wariancję estymatora DEA oraz wyprowadzić postaci odpowiednich asymptotycznych przedziałów ufności. Na koniec wykonano badania symulacyjne, mające na celu sprawdzenie małopróbkowych własności estymatora DEA oraz zastosowano przedstawioną metodologię do analizy poziomu efektywności rzeczywistych obiektów produkcyjnych z polskiego sektora energetycznego.

Słowa kluczowe: metoda DEA, efektywność techniczna, rozkład asymptotyczny, nieparametryczny model graniczny.

ESTIMATION OF THE TECHNICAL EFFICIENCY MEASURE WITHIN THE DEA METHOD

Summary

In the paper we present the proposition of the description of uncertainty related to the technical efficiency measure received using the DEA method. A nonparametric, statistical model functioning within the DEA method is introduced and the DEA estimator of the technical efficiency measure is defined. For a single input – single output variable case, the value of the technical efficiency measure and frontier production function are related. Thus a corresponding DEA frontier estimator for a given point is defined. Next, forms of the asymptotic distributions of the DEA estimators are presented. It allowed us to derive the asymptotic bias and variance of the DEA estimator and to construct asymptotic confidence intervals. In the last part, finite sample performance of the DEA estimator is investigated via a simulation study. We also illustrate the DEA estimation procedure using real data, which come from the Polish Energy Sector.

Key words: DEA method, technical efficiency, asymptotic distribution, nonparametric Frontier Model.

SABINA DENKOWSKA, KAMIL FIJOREK, MARCIN SALAMAGA, ANDRZEJ SOKOŁOWSKI

EMPIRYCZNA OCENA MOCY TESTÓW DLA WIELU WARIANCJI

1. WPROWADZENIE

Testy dla wielu wariancji nie są zbyt popularne wśród testów statystycznych. Może dlatego, że rzadko wykorzystywane są jako procedura docelowa. Raczej służą do weryfikowania założenia o równości wariancji wymaganego przez inne procedury – przede wszystkim analizę wariancji. Mogą być też wykorzystywane do oceny jednorodności źródeł danych, które chcemy połączyć dla uzyskania lepszych ocen wariancji.

Testy dla wielu wariancji też zazwyczaj same wymagają spełnienia pewnych założeń. Najczęściej chodzi tu o założenie normalności rozkładu w populacjach lub przynajmniej rozkładu tego samego typu.

Oczywiście, zagadnienie porównywania rzeczywistego prawdopodobieństwa błędu pierwszego rodzaju oraz mocy testów dla wielu wariancji było już podejmowane w literaturze statystycznej. Chyba najbardziej wyczerpującym doświadczeniem były symulacje, których wyniki przedstawili Conover, Johnson, Johnson (1981)¹. Poddali oni analizie m.in. następujące testy: Neymana i Pearsona (1931), Bartletta (1937), Cochraha (1941), Scheffego (1959), Levene'a (1960), Browna i Forsythe'a (1974).

Zasadniczym celem artykułu jest zbadanie zachowania się wybranych testów dla wielu wariancji w warunkach prawdziwości i nieprawdziwości założenia o normalności rozkładu, małych i dużych prób, prób równolicznych i prób o zróżnicowanych liczebnościach. Badania te przeprowadzono metodami symulacyjnymi.

Z bogatego zbioru testów dla wielu wariancji do analizy wybrano testy wyróżnione w badaniach Conovera, Johnsona, Johnsona z 1981 r. (test Levene'a i test Flingera-Killeena), test Browna-Forsythe'a, który jest często stosowany ze względu na dostępność w komputerowych pakietach statystycznych oraz relatywnie nową propozycję O'Briena z 1981 r.²

2. WYBRANE TESTY JEDNORODNOŚCI WARIANCJI

Założenie jednorodności (równości) wariancji w obrębie grup jest jednym z wymogów ważnych procedur wnioskowania statystycznego. Spełnienie tego założenia jest wymagane m.in. w przypadku stosowania jednoczynnikowej analizy wariancji ANOVA.

¹ Por. [2].

² Por. [1], s. 1.

Do bardziej znanych testów jednorodności wariancji obok testów Hartleya, Cochрана i Bartletta należy procedura Levene'a³.

W teście homogeniczności wariancji weryfikujemy następujące hipotezy statystyczne:

$$H_0: \sigma_1^2 = \sigma_2^2 = \sigma_3^2 = \dots = \sigma_k^2$$

$$H_1: \sigma_i^2 \neq \sigma_j^2 \text{ przynajmniej dla jednej pary } (i, j).$$

Założenia procedury Levene'a (1960) są następujące:

1. Cecha X ma w i -tej populacji rozkład normalny $N(\mu_i, \sigma_i^2)$, gdzie $i = 1, 2, \dots, k$,
2. Wartości przeciętne μ_i w k populacjach są nieznane.

W teście Levene'a dla każdej obserwacji X_{ij} wyznaczane jest odchylenie absolutne względem wartości średniej $\bar{X}_i$ w i -tej grupie:

$$Z_{ij} = |X_{ij} - \bar{X}_i|, j = 1, 2, \dots, n_i, i = 1, 2, \dots, k. \quad (1)$$

Statystyka testowa zaproponowana przez Levene'a ma postać⁴:

$$W = \frac{N-1}{k-1} \frac{\sum_{i=1}^k n_i (\bar{Z}_i - \bar{Z})^2}{\sum_{i=1}^k \sum_{j=1}^{n_i} (Z_{ij} - \bar{Z}_i)^2}, \quad (2)$$

przy czym: $N = \sum_{i=1}^k n_i$, natomiast $\bar{Z}_i$ oraz $\bar{Z}$ oznaczają odpowiednio średnią wewnątrzgrupową i średnią międzygrupową. Przy założeniu prawdziwości hipotezy zerowej statystyka W ma rozkład Fishera o $k-1$ oraz $N-k$ stopniach swobody. Jeżeli na poziomie istotności α empiryczna wartość statystyki W przekracza wartość krytyczną $F_{\alpha; k-1; N-k}$, to hipotezę zerową odrzucamy. Test Levene'a został potem zmodyfikowany przez Browna i Forsythe'a (1974) poprzez zastąpienie średniej wewnątrzgrupowej $\bar{X}_i$ w przekształceniu (1) medianą lub średnią uciętą. Autorzy modyfikacji testu Levene'a zalecają przede wszystkim stosowanie procedury z medianą grupową, gdyż zapewnia ona, ich zdaniem, test o dość wysokiej mocy, który jest odporny na brak normalności w rozkładzie danych.

Kolejną procedurą zaproponowaną do testowania jednorodności wariancji w k -populacjach jest test Flignera-Killeena (1974). Ta procedura wymaga utworzenia rankingu bezwzględnych wartości cechy X_{ij} . Uporządkowanym niemalejąco wartościom $|X_{ij}|$ przypisuje się następnie wskaźniki a_{Ni} obliczane następująco⁵:

$$a_{Ni} = \Phi^{-1} \left(\frac{1}{2} + \frac{i}{2(N+1)} \right). \quad (3)$$

³ Por. [5], s. 278-292.

⁴ Por. [7], s. 49.

⁵ Por. [6], s. 3.

Warto zaznaczyć, że Conover, Johnson, Johnson (1981) zaproponowali rangowanie wartości bezwzględnych cechy X_{ij} po uprzednim odjęciu median grupowych.

Statystyka testowa w procedurze Flignera-Killeena (1974) ma postać:

$$\chi^2 = \sum_{i=1}^k \frac{n_i(\bar{A}_i - \bar{a})^2}{V^2}, \quad (4)$$

przy czym $\bar{A}_i$ jest średnią wewnątrzgrupową, $\bar{a}$ jest średnią międzygrupową oraz V^2 jest wariancją międzygrupową. Statystyka (4) ma asymptotyczny rozkład chi-kwadrat o $k - 1$ stopniach swobody. Jeżeli na poziomie istotności α empiryczna wartość statystyki χ^2 przekracza wartość krytyczną $\chi^2_{\alpha, k-1}$, to hipotezę zerową (głoszącą równość wariancji w k populacjach) odrzucamy.

Nieco inne podejście w testowaniu jednorodności wariancji zaproponował O'Brien (1979, 1981). Procedura O'Briena wykorzystuje jednoczynnikową analizę wariancji ANOVA, która jest stosowana dla przekształconych wartości cechy X_{ij} zgodnie ze wzorem⁶:

$$Z_{ij} = \frac{n_i(n_i - 1,5)(X_{ij} - \bar{X}_i)^2 - 0,5 \sum_{j=1}^{n_i} (X_{ij} - \bar{X}_i)^2}{(n_i - 1)(n_i - 2)}, \quad (5)$$

$$j = 1, 2, \dots, n_i, i = 1, 2, \dots, k,$$

przy czym:

n_i – liczebność i -tej grupy,

$\bar{X}_i$ – wartości średniej w i -tej grupie.

Do weryfikacji hipotezy o równości wariancji stosowany jest potem test F Fishera:

$$F = \frac{(N - k) \sum_{i=1}^k n_i (\bar{Z}_i - \bar{Z})^2}{(k - 1) \sum_{i=1}^k \sum_{j=1}^{n_i} (Z_{ij} - \bar{Z}_i)^2}, \quad (6)$$

przy czym: $N = \sum_{i=1}^k n_i$, natomiast $\bar{Z}_i$ oraz $\bar{Z}$ oznaczają odpowiednio średnią wewnątrzgrupową i międzygrupową.

Statystyka (6) ma rozkład F o $k - 1$ i $N - k$ stopniach swobody. Jeżeli empiryczna wartość statystyki (6) przekracza wartość krytyczną $F_{\alpha; k-1; N-k}$, to hipotezę zerową o równości wariancji w k populacjach odrzucamy.

⁶ Por. [1], s. 3.

3. METODOLOGIA BADAŃ SYMULACYJNYCH

Celem przeprowadzenia badań symulacyjnych było określenie zachowania się wybranych testów jednorodności wariancji rozumianego jako zdolność (lub jej brak) do utrzymania nominalnego prawdopodobieństwa popełnienia błędu I rodzaju oraz zdolność do uzyskiwania wysokiego prawdopodobieństwa odrzucenia fałszywej hipotezy zerowej (moc testu) w warunkach zmieniających się liczebności prób oraz zmieniających się konfiguracji wielkości wariancji w poszczególnych próbach. W badaniu rozważono cztery testy jednorodności wariancji: Levene'a, Browna-Forsythe'a, Fligner-Killeena oraz O'Brien'a.

Wykorzystanie symulacji komputerowych jest podyktowane praktyczną niemożliwością analitycznego wyznaczenia krzywych mocy dla wybranych testów jednorodności wariancji.

W przeprowadzonym badaniu rozważono trzy rozkłady prawdopodobieństwa (mechanizmy generujące próby):

- rozkład normalny,
- rozkład logarytmiczno-normalny,
- rozkład jednostajny.

W przypadku testów jednorodności wariancji wielkość wartości oczekiwanej nie ma wpływu na wyniki analizy. W badaniach przyjęto arbitralnie dla rozkładu logarytmiczno-normalnego oraz normalnego wartość oczekiwaną równą 10. W przypadku rozkładu jednostajnego arbitralnie przyjęto wartość 0 jako dolną granicę rozkładu. Uzyskanie pożądanych wartości odchylenia standardowego było możliwe poprzez odpowiednie sterowanie górną granicą przedziału, na którym jest określony rozkład jednostajny.

W przeprowadzonej symulacji założono, że obserwacje pochodzą z czterech populacji. Ponadto zbudowano trzy mechanizmy określania wielkości odchylenia standardowego w każdej z czterech populacji (oznaczonych dolnymi subskryptami):

- $\sigma_1 = \sigma_2 = \sigma_3 = 1$, $\sigma_4 = 1,0 + 3^*$ (współczynnik wzrostu odchylenia standardowego),
- $\sigma_1 = \sigma_2 = 1$, $\sigma_3 = \sigma_4 = 1,0 + 3^*$ (współczynnik wzrostu odchylenia standardowego),
- $\sigma_i = 1,0 + (i - 1)^*$ (współczynnik wzrostu odchylenia standardowego), $i = 1, 2, 3, 4$.

Współczynnik wzrostu odchylenia standardowego przyjmuje wartości od 0 do 0,35 z krokiem 0,025.

W badaniu założono następujące sześć możliwych konfiguracji liczebności prób: (10, 10, 10, 10), (40, 40, 40, 40), (5, 10, 15, 20), (20, 15, 10, 5), (30, 40, 50, 60), (60, 50, 40, 30). Wykorzystanie zarówno prób równolicznych, jak i nierównolicznych jest podyktowane chęcią zbadania wpływu tego czynnika na zachowanie się rozważanych testów statystycznych.

Ponadto uwzględnienie liczebności prób (5, 10, 15, 20), (20, 15, 10, 5), (30, 40, 50, 60), (60, 50, 40, 30) pozwala zbadać łączny wpływ wielkości odchylenia standardowego w danej próbie oraz liczebności tej próby na rozważane testy jednorodności wariancji.

Wszystkie symulacje przeprowadzono w środowisku statystycznym R (www.r-project.org) w oparciu o autorskie skrypty obliczeniowe.

4. WYNIKI BADAŃ

4.1. KONTROLA BŁĘDU PIERWSZEGO RODZAJU

W sytuacji prawdziwości hipotezy zerowej badania symulacyjne pozwoliły oszacować rozmiary rozpatrywanych testów⁷ jednorodności wariancji. W celu oszacowania prawdopodobieństwa popełnienia błędu I rodzaju dziesięć tysięcy razy generowano cztery próby o różnej liczności z wybranych rozkładów i za każdym razem badano, która z procedur nie rozpozna stanu faktycznego, czyli równości wariancji we wszystkich grupach. W tabelach 1-3 przedstawione są wyniki otrzymane w zależności od liczebności podgrup.

Ważnym założeniem większości testów jednorodności wariancji jest założenie mówiące o normalności populacji, z których losujemy próby. W prowadzonych badaniach symulacyjnych szacowano prawdopodobieństwo błędu I rodzaju w sytuacji, gdy założenie o normalności jest spełnione (patrz: tabela 1), ale również badano wrażliwość tych procedur na niespełnienie tego założenia. Tabela 2 prezentuje wyniki dla prób generowanych z rozkładu lognormalnego, a tabela 3 z rozkładu jednostajnego.

Tabela 1

Oceny prawdopodobieństw odrzucenia prawdziwej hipotezy zerowej – równe odchylenia standardowe, rozkłady normalne

Liczebności grup				TEST			
I	II	III	IV	Levene'a	B-F	F-K	O'Briena
10	10	10	10	0,0694	0,0346	0,0358	0,0389
40	40	40	40	0,0547	0,0430	0,0424	0,0453
5	10	15	20	0,0814	0,0368	0,0379	0,0608
30	40	50	60	0,0616	0,0501	0,0511	0,0524

Źródło: opracowanie własne.

W tabeli 1 przedstawiono wyniki uzyskane, gdy próby były losowane z populacji normalnych o równych wariancjach. Najgorzej wypadł test Levene'a, który dla wszystkich rozpatrywanych wariantów liczebności przekroczył przyjęty na wstępie badań poziom istotności 0,05. Szczególnie wysokie prawdopodobieństwo błędnego odrzucenia prawdziwej hipotezy zerowej można zauważyć dla małych, nierównolicznych prób. Procedura ta nie wypada dobrze również w przypadku dużych, ale nierównolicznych próbek. Najlepiej w tej części badań wypadają procedury Browna-Forsythe'a oraz Flignera-Killeena. Obie zapewniają kontrolę błędu I rodzaju na poziomie 0,05, nieznacznie przekraczając to prawdopodobieństwo w przypadku dużych, nierównolicznych prób. Pomiędzy procedurami Browna-Forsythe'a oraz Flignera-Killeena uplasowała się

⁷ Patrz [4], s. 23.

procedura O'Briena, która zapewnia kontrolę błędu I rodzaju dla prób równolicznych, również małych. W przypadku próbek o różnej liczności oszacowane prawdopodobieństwa, mimo iż większe od 0,05, są jednak znacznie niższe od prawdopodobieństw otrzymanych przy zastosowaniu testu Levene'a.

Tabele 2 i 3 przedstawiają wyniki w przypadku, gdy niespełnione są założenia o normalności rozkładów.

W tabeli 2 przedstawione są wyniki symulacji prowadzonych dla prób generowanych z rozkładu lognormalnego. Zaskakują bardzo złe wyniki testu Levene'a. Oszacowane prawdopodobieństwa popełnienia błędu I rodzaju przekraczają 0,3! Okazuje się, że zastosowanie tego testu do prób z rozkładu lognormalnego powoduje odrzucenie co najmniej 30% prawdziwych hipotez zerowych. Najbardziej odporny okazał się test Browna-Forsythe'a, który jedynie w sytuacji małych, nierównolicznych prób przekroczył założony na wstępie poziom istotności 0,05. W miarę odporny okazał się również test O'Briana, który kontrolę prawdopodobieństwa odrzucenia hipotezy prawdziwej zapewniał w przypadku dużych prób. Niestety, prawdopodobieństwo błędu I rodzaju dwukrotnie przekracza dopuszczalny poziom istotności w przypadku zastosowania tego testu do małych nierównolicznych próbek. Test Flignera-Killeena we wszystkich rozpatrywanych wariantach liczebności ponad dwukrotnie przekraczał α .

Tabela 2

Oceny prawdopodobieństw odrzucenia prawdziwej hipotezy zerowej – równe odchylenia standardowe, rozkłady lognormalne

Liczebności grup				TEST			
I	II	III	IV	Levene'a	B-F	F-K	O'Briena
10	10	10	10	0,3438	0,039	0,1205	0,0584
40	40	40	40	0,3013	0,0428	0,1533	0,0344
5	10	15	20	0,3328	0,057	0,1175	0,1044
30	40	50	60	0,3059	0,0485	0,1654	0,0405

Źródło: opracowanie własne.

Tabela 3 przedstawia wyniki uzyskane dla prób losowanych z rozkładu jednostajnego. I w tym przypadku najgorzej wypadła procedura Levene'a, która dla wszystkich czterech wariantów liczebności dała oszacowania przekraczające 0,05. Uzyskane przez tę procedurę wyniki są zbliżone do rezultatów uzyskanych dla rozkładu normalnego (por. tabela 1) i potwierdza się spostrzeżenie, że procedura ta nie wypada dobrze w przypadku małych, nierównolicznych prób, jak również w przypadku dużych, ale nierównolicznych prób. W sytuacji rozkładu prostokątnego zarówno procedura Browna-Forsythe'a, jak i Flignera-Killeena wypadły w badaniach bardzo dobrze, uzyskując prawdopodobieństwa błędnej decyzji poniżej 0,05. Dla testu O'Briena, podobnie jak w przypadku rozkładu normalnego, otrzymano nieznaczne przeszacowanie prawdopodobieństwa błędu I rodzaju w przypadku prób nierównolicznych.

Tabela 3

Oceny prawdopodobieństw odrzucenia prawdziwej hipotezy zerowej – równe odchylenia standardowe, rozkłady jednostajne

Liczebności grup				TEST			
I	II	III	IV	Levene'a	B-F	F-K	O'Briena
10	10	10	10	0,0709	0,025	0,0232	0,0411
40	40	40	40	0,0562	0,033	0,0332	0,0471
5	10	15	20	0,097	0,0374	0,0309	0,0665
30	40	50	60	0,0607	0,0389	0,0379	0,0523

Źródło: opracowanie własne.

W prowadzonych badaniach rozważano trzy rozkłady, z których generowano próby losowe. Okazało się, że wyniki jakie uzyskano dla rozkładu jednostajnego tylko nieznacznie różnią się od wyników otrzymanych dla rozkładu normalnego, który stanowi kluczowe założenie większości procedur służących do badania jednorodności wariancji. Natomiast bardzo istotne różnice wystąpiły w przypadku zastosowania testów do prób wygenerowanych z rozkładu lognormalnego. W przypadku gdy $\sigma \geq 1$ rozkład lognormalny charakteryzuje się silną prawostronną asymetrią i wydaje się, że właśnie to odejście od symetrii jest powodem tak znacznego pogorszenia oszacowań błędów I rodzaju procedur Levene'a i Flignera-Killeena.

Podsumowując część badań dotyczącą kontroli błędu I rodzaju możemy stwierdzić, że pod tym względem zdecydowanie najgorzej wypadła procedura Levene'a. Okazała się „nieodporna” w przypadku nierównych lub nielicznych prób, nawet tych losowanych z rozkładu normalnego. W przypadku rozkładu asymetrycznego jakim jest rozkład lognormalny, wyniki tej procedury pogarszają się tak drastycznie, że jej stosowanie wydaje się wręcz niedopuszczalne. Najlepiej spośród rozważanych testów wypadł test Browna-Forsythe'a. Test ten okazał się znacznie bardziej od innych procedur odporny na nierównoliczność prób, czy też odchylenia od normalności. Stosunkowo dobrze wypadł on również w przypadku silnej asymetrii rozkładu lognormalnego, przekraczając zaledwie o 0,007 założony poziom istotności dla prób małych i nierównolicznych. Z pozostałych testów ciekawie wypada test O'Briena, dla którego oszacowane prawdopodobieństwa testowe oscylują koło 0,05, nieznacznie przekraczając tę wartość dla prób nierównolicznych. W przypadku silnej asymetrii rozkładu lognormalnego test ten okazał się nieodporny tylko w przypadku małych prób, w przeciwieństwie do testu Flignera-Killeena, dla którego wszystkie oszacowane prawdopodobieństwa odrzucenia prawdziwej hipotezy zerowej okazały się ponad dwukrotnie większe od przyjętego poziomu istotności.

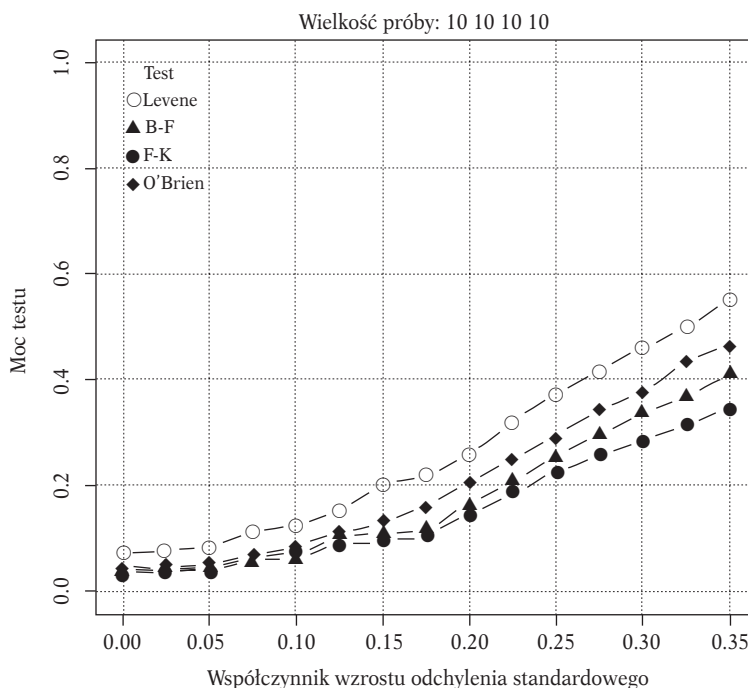
4.2. BADANIE MOCY PROCEDUR

Mocą testu nazywamy prawdopodobieństwo podjęcia słusznej decyzji polegającej na odrzuceniu fałszywej hipotezy zerowej. Druga część badań miała na celu porów-

nanie mocy wybranych testów służących do badania równości wielu wariancji. W tym przypadku tysiąc razy generowano cztery próby o różnej liczebności z wybranych rozkładów i za każdym razem badano, która z procedur rozpozna stan faktyczny, czyli brak równości wariancji we wszystkich grupach. Dla każdej z procedur szacowano prawdopodobieństwa wykrycia niejednorodności wariancji.

Istotnym założeniem wielu testów jednorodności wariancji jest założenie dotyczące normalności rozkładów, z których pobierane są próby. Niestety, w praktyce często założenie to nie jest spełnione, a mimo to procedury są stosowane. W badaniach sprawdzano więc moc procedur nie tylko w sytuacji, gdy próby są generowane z rozkładów normalnych. Generowano również próby z rozkładu jednostajnego i lognormalnego. Tak więc symulacje miały na celu również ocenę odporności wybranych testów w przypadku niespełnienia założenia o normalności.

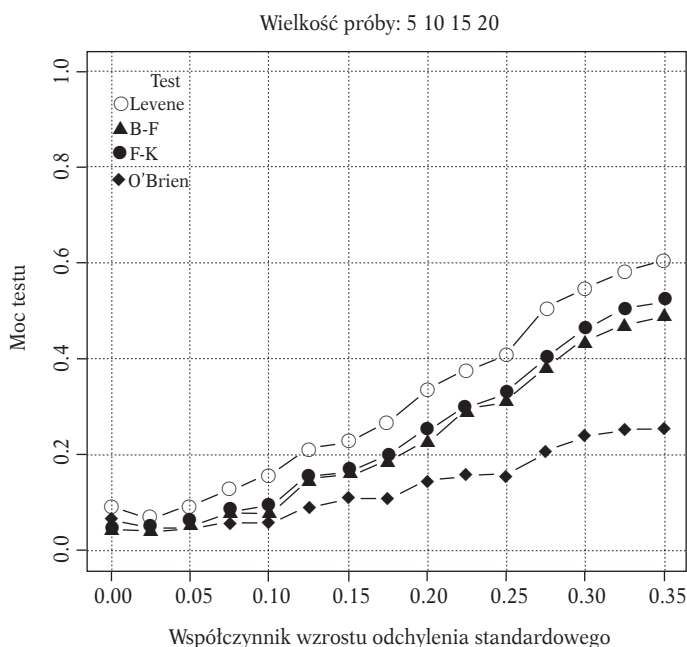
Na wstępie, w celu porównania mocy procedur w sytuacji spełnionych założeń, eksperymenty przeprowadzono dla prób generowanych z rozkładów normalnych. Wszystkie cztery rozpatrywane testy osiągały porównywalne rezultaty w przypadku dużych prób. Oczywiście, prawdopodobieństwa rozpoznania niejednorodności wariancji rosły wraz ze wzrostem różnic pomiędzy wariancjami. Jednak wykrywalność niejednorodności wariancji okazała się wyraźnie lepsza, gdy rosła wariancja tylko w jednej lub dwóch grupach. Dla obu tych sytuacji oszacowane prawdopodobieństwa różniły się nieznacznie.



Rysunek 1. Oceny mocy testów dla prób z rozkładu normalnego. Odchylenia standardowe wynoszą odpowiednio: $\sigma_1 = \sigma_2 = \sigma_3 = 1$ i $\sigma_4 = 1 + 3^*$ (współczynnik wzrostu odchylenia standardowego)

Wyraźne różnice pomiędzy testami dało się zaobserwować w wykrywalności niejednorodności wariancji w przypadku małych prób. W przypadku prób równolicznych 10-elementowych zdecydowanie najlepszym rozpoznaniem niejednorodności wykazała się procedura Levene'a. Procedura ta jednak nie zapewnia kontroli błędu I rodzaju na założonym poziomie istotności, co stanowi jej poważny mankament. Dla pozostałych trzech procedur różnice w oszacowanych prawdopodobieństwach były nieznaczne, ale wyraźnie gorsze od prawdopodobieństw otrzymanych dla procedury Levene'a. Tylko w sytuacji, gdy wariancja rosła w jednej grupie (patrz: rys. 1) różnice między testami były wyraźniejsze i spośród testów kontrolujących prawdopodobieństwo błędu I rodzaju najlepiej wypadł test O'Brien, trochę gorzej test Browna-Forsythe'a, a najgorzej test Flignera-Killeena. Oczywiście, różnice pomiędzy testami rosły wraz ze wzrostem wariancji w grupie.

W przypadku małych nierównolicznych prób okazało się, iż istotny wpływ na rozpoznanie niejednorodności wariancji przez różne testy ma fakt, czy mniejszym próbom w eksperymencie odpowiadały mniejsze czy większe wariancje. I tak np. procedura O'Brien była najlepsza, gdy większym liczebnościom odpowiadały mniejsze wariancje, podczas gdy w sytuacji odwrotnej wypadała ona najgorzej w rankingu testów jednorodności wariancji (patrz: rys. 2).



Rysunek 2. Oceny mocy testów dla prób z rozkładu normalnego. Odchylenia standardowe w poszczególnych grupach wynoszą odpowiednio: $\sigma_1 = \sigma_2 = 1$ i $\sigma_3 = \sigma_4 = 1 + 3^*$ (współczynnik wzrostu odchylenia standardowego)

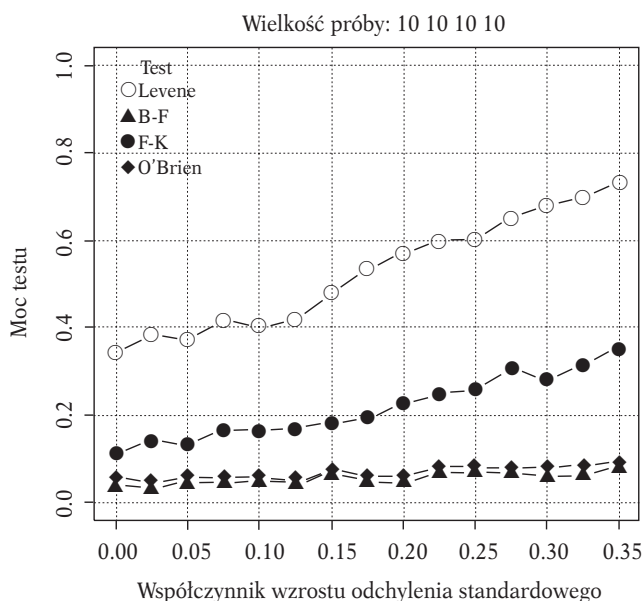
Podsumowując, eksperymenty symulacyjne przeprowadzone na próbach generowanych z rozkładów normalnych, można stwierdzić, iż w przypadku dużych prób wszystkie badane procedury miały zbliżone wyniki rozpoznania niejednorodności podgrup. W przypadku małych prób sytuacja była bardziej skomplikowana. Najlepsze prawdopodobieństwa rozpoznania niejednorodności miał test Levene'a, który jednak nie zapewniał kontroli błędu I rodzaju na założonym poziomie istotności.

W przypadku procedury O'Briena wyniki, w porównaniu do pozostałych testów, wahały się od bardzo dobrych do bardzo słabych, a zależało to od tego, czy liczniejszym próbom odpowiadały mniejsze czy większe wariancje. Wydaje się, że taka „niestabilność” dyskwalifikuje procedurę O'Briena w badaniach praktycznych w przypadku małych prób.

Dobre, stabilne wyniki miała procedura Browna-Forsythe'a, która w przeciwieństwie do procedury Levene'a kontrolowała prawdopodobieństwo odrzucenia prawdziwej hipotezy zerowej na przyjętym w eksperymencie poziomie istotności 0,05.

Ponieważ w badaniach praktycznych nie mamy możliwości rozpoznania, czy mniejszym próbom odpowiada większa czy mniejsza wariancja (chyba, że oprzemy się na obserwacjach dla prób) wydaje się więc, że najbezpieczniej w przypadku małych, nierównolicznych prób decyzję o niejednorodności wariancji przeprowadzać za pomocą np. testu Browna-Forsythe'a.

Kolejne eksperymenty symulacyjne polegały na losowaniu prób z rozkładu lognormalnego i jednostajnego o zadanych parametrach i badaniu wrażliwości testów na odejście od założenia o normalności.

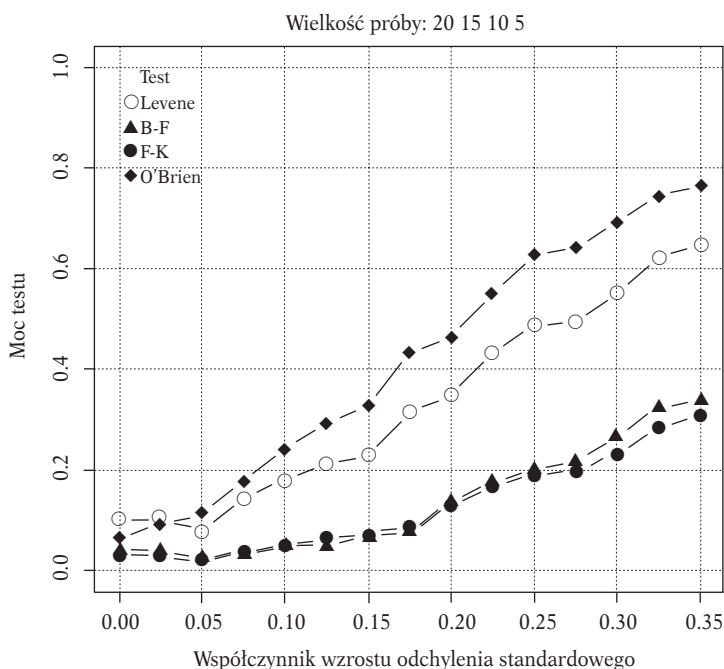


Rysunek 3. Oceny mocy testów dla prób z rozkładu lognormalnego. Odchylenia standardowe wynoszą odpowiednio dla i -tej grupy: $\sigma_4 = 1 + (i - 1) \cdot$ (współczynnik wzrostu odchylenia standardowego)

Interesująco, ale zarazem niepokojąco wypadły badania, w których próby były generowane z rozkładu lognormalnego. Zdecydowanie najlepiej niejednorodność wariancji wykrywała procedura Levene'a (por. rys. 3). W przypadku małych, nierównolicznych prób z rozkładu lognormalnego okazało się, że istotna jest zależność pomiędzy liczebnością próby i wielkością wariancji. I tak, w przypadku gdy liczniejszej próbie odpowiada większa wariancja (bez znaczenia jest fakt, czy zmiany wariancji zachodziły w jednej, dwóch czy w trzech grupach) procedura Flignera-Killeena wydaje się jedynym rozwiązaniem. Natomiast gdy mniejszej próbie odpowiada większa wariancja, wyraźnie poprawia się moc pozostałych procedur (Brown-Forsythe'a i O'Briena).

W przypadku dużych prób generowanych z rozkładów jednostajnych wszystkie testy uzyskały prawdopodobieństwa rozpoznania niejednorodności wariancji co najmniej tak dobre, jak w analogicznym eksperymencie z próbami generowanymi z rozkładu normalnego. Wśród rozważanych testów najlepiej wypadła procedura O'Briena.

Procedura O'Briena również w przypadku małych prób dała najlepsze oszacowania prawdopodobieństw rozpoznania niejednorodności wariancji. Szczególnie duży kontrast w stosunku do pozostałych procedur pojawiał się wtedy, gdy próby były równoliczne lub gdy mniejszej próbie odpowiadała większa wariancja (rys. 4).



Rysunek 4. Oceny mocy testów dla prób z rozkładu jednostajnego, w przypadku gdy mniejszej próbie odpowiada większa wariancja. Odchylenia standardowe wynoszą dla i -tej grupy odpowiednio:

$$\sigma_i = 1 + (i - 1) * (\text{współczynnik wzrostu odchylenia standardowego})$$

W najgorszym przypadku, gdy najmniejszej próbie odpowiadała najmniejsza wariancja, procedura O'Briena miała wyniki zbliżone do wyników pozostałych procedur: Browna-Forsythe'a i Flignera-Killeena. Najlepiej w tym eksperymencie wypadła procedura Levene'a, należy jednak wspomnieć, że również w tym eksperymencie test ten dał prawie dwukrotne przeszacowanie przyjętego na wstępie poziomu istotności.

5. PODSUMOWANIE

Przeprowadzone badania symulacyjne potwierdziły, że nie ma jednego najlepszego, czy też „uniwersalnego” testu jednorodności wariancji. Zalecenia autorów, wynikające z przeprowadzonych badań, odnośnie wyboru najwłaściwszego testu jednorodności w zależności od rozważanych czynników, takich jak m.in. rozkład, czy liczebności prób, przedstawione są w tabelach 4-5. W kolumnie „zalecenia” wymienione są procedury w zalecanej kolejności stosowania.

Tabela 4

Zalecenia odnośnie wyboru testu jednorodności wariancji w przypadku losowania prób z rozkładów co najmniej zbliżonych do rozkładu normalnego lub do rozkładu jednostajnego

Liczebności prób		Zalecenia (rozkłady zbliżone do rozkładu normalnego)	Zalecenia (rozkłady zbliżone do rozkładu jednostajnego)
Małe	Równoliczne	1. Test O'Briena 2. Test Browna-Forsythe'a 3. Test Flignera-Killeena 4. Test Levene'a ($p = 0,0694$)	1. Test O'Briena 2. Test Levene'a ($p = 0,0709$) 3. Test Browna-Forsythe'a 4. Test Flignera-Killeena
	Nierównoliczne	1. Test Browna-Forsythe'a 2. Test Flignera-Killeena 3. Test Levene'a ($p = 0,0814$) 4. Test O'Briena ($p = 0,0608$)	1. Test O'Briena ($p = 0,0665$) 2. Test Levene'a ($p = 0,097$) 3. Test Browna-Forsythe'a 4. Test Flignera-Killeena
Duże	Równoliczne	1. Test O'Briena 2. Test Levene'a ($p = 0,0547$) 3. Test Browna-Forsythe'a 4. Test Flignera-Killeena	1. Test O'Briena 2. Test Flignera-Killeena 3. Test Levene'a ($p = 0,0562$) 4. Test Browna-Forsythe'a
	Nierównoliczne	1. Test O'Briena ($p = 0,0524$) 2. Test Levene'a ($p = 0,0616$) 3. Test Browna-Forsythe'a ($p = 0,0501$) 4. Test Flignera-Killeena ($p = 0,0511$)	1. Test O'Briena ($p = 0,0523$) 2. Test Flignera-Killeena 3. Test Levene'a ($p = 0,0607$) 4. Test Browna-Forsythe'a

Źródło: opracowanie własne (w nawiasach obok nazw procedur podano oszacowane dla tych procedur prawdopodobieństwa popełnienia błędu I rodzaju większe od 0,05).

Część z zalecanych testów nie zapewnia kontroli błędu I rodzaju. Ponieważ dla rozpatrywanych procedur różnice w oszacowaniach prawdopodobieństw błędu I rodzaju były znaczne w zależności od rodzaju procedury, czy też rozważanej sytuacji badawczej, w nawiasie obok nazwy procedury podano oszacowane prawdopodobieństwa popełnienia błędu I rodzaju, wtedy gdy jest ono większe od 0,05. Kolejność zalecanych testów zależała od kilku czynników. Oczywiście, autorzy starali się wyróżnić testy o największej mocy, zapewniające kontrolę błędu I rodzaju. W niektórych sytuacjach badawczych wybór najlepszych procedur był szczególnie trudny, gdy np. procedury dobrze rozpoznające niejednorodność wariancji, miały zbyt wysokie prawdopodobieństwa błędu I rodzaju. Prawdopodobieństwa popełnienia błędu I rodzaju zostały więc podane, by badacz korzystający z tego opracowania miał orientację, na jaki błąd I rodzaju się naraża dokonując wyboru procedury i ewentualnie rozważył, czy jest to dopuszczalne w prowadzonych przez niego badaniach. Jeśli nie dopuszcza takiego poziomu błędu I rodzaju, może rozważyć zastosowanie testu na „odpowiednio” niższym poziomie istotności lub też zastosować kolejny test z listy zaleceń.

Tabela 5

Zalecenia odnośnie wyboru testu jednorodności wariancji w przypadku losowania prób z rozkładów o wyraźnej asymetrii umiarkowanej

Liczebności prób		Zalecenia
Małe	Równoliczne	Test Flignera-Killeena ($p = 0,1205$)
	Nierównoliczne	Żaden z testów nie zasługuje na rekomendację, ale „najmniejszym złem” wydaje się test Flignera-Killeena ($p = 0,1175$).
Duże	Równoliczne i nierównoliczne	1. Test Flignera-Killeena ($p = 0,1700$) 2. Test Browna-Forsythe’a

Źródło: opracowanie własne (w nawiasach obok nazw procedur podano oszacowane dla tych procedur prawdopodobieństwa popełnienia błędu I rodzaju większe od 0,05).

W tabeli 4 podane są zalecenia w sytuacji, gdy próby są losowane z rozkładów normalnych (lub rozkładów co najmniej zbliżonych do normalnych) oraz dla prób losowanych z rozkładów lognormalnych. W badaniach rozważano odchylenia standardowe, dla których rozkład lognormalny charakteryzuje się wyraźną asymetrią prawostronną. Wydaje się, że zalecenia z tej części badań można odnieść nie tylko do rozkładu lognormalnego, ale do szerszej klasy rozkładów umiarkowanie asymetrycznych.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] Abdi H., [2007], *O'Brien test for homogeneity of variance*, [w:] N.J. Salkind (red.), *Encyclopedia of Measurement and Statistics*, Thousand Oaks (CA), Sage, pp. 701-704.

- [2] Conover W.J., Johnson M.E., Johnson M.M., [1981], *A Comparative Study of Tests for Homogeneity of Variances, with Applications to the Outer Continental Shelf Bidding Data*, Technometrics, Vol. 23, No. 4, s. 351-361.
- [3] Domański Cz., [1986], *Teoretyczne podstawy testów nieparametrycznych i ich zastosowanie w naukach ekonomiczno-społecznych*, Acta Universitatis Lodziensis, Wyd. UŁ, s. 233-240.
- [4] Domański Cz., [1990], *Testy statystyczne*, PWE.
- [5] Domański Cz., Pruska K., [2000], *Nieklasyczne metody statystyczne*, PWE.
- [6] Flinger M.A., Killeen T.J., [1976], *Distribution-Free Two-Sample Tests for Scale*, „Journal of the American Statistical Association”, 71, 210-213.
- [7] Lehmann E.L., (1959), *Testing Statistical Hypothesis*, New York, John Wiley.
- [8] Levene, H., [1960], *In Contributions to Probability and Statistics: Essays in Honor of Harold Hotelling*, I. Olkin et al. (red.), Stanford University Press, s. 278-292.
- [9] Xu-Feng N., [2004], *Statistical Procedures for Testing Homogeneity of Water Quality Parameters*, Department of Statistics Florida State University Tallahassee.
- [10] Zieliński W., [1999], *Wybrane testy statystyczne*, Fundacja „Rozwój SGGW”, Warszawa.

Praca wpłynęła do redakcji w kwietniu 2009 r.

EMPIRYCZNA OCENA MOCY TESTÓW DLA WIELU WARIANCJI

Streszczenie

Testy dla wielu wariancji stosuje się zwykle do weryfikowania założenia o równości wariancji wymaganego przez inne procedury – przede wszystkim analizę wariancji. Mogą być też wykorzystywane do oceny jednorodności źródeł danych, które chcemy połączyć dla uzyskania lepszych ocen wariancji.

Testy dla wielu wariancji też zazwyczaj same wymagają spełnienia pewnych założeń. Najczęściej chodzi tu o założenie normalności rozkładu w populacjach lub przynajmniej rozkładu tego samego typu.

Celem artykułu jest zbadanie zachowania się wybranych testów dla wielu wariancji w warunkach prawdziwości i nieprawdziwości założenia o normalności rozkładu, małych i dużych prób, prób równolicznych i prób o zróżnicowanych liczebnościach. Badania te przeprowadzono metodami symulacyjnymi. W badaniu rozważono cztery testy jednorodności wariancji: Levene’a, Browna-Forsythe’a, Flignera-Killeena oraz O’Brien’a.

Słowa kluczowe: testy dla wielu wariancji, symulacje Monte Carlo, test Levene’a, test Browna-Forsythe’a, test Flignera-Killeena, test O’Brien’a.

A COMPARATIVE STUDY OF TESTS POWER FOR HOMOGENEITY OF VARIANCES

Summary

Some statistical tests, for example the analysis of variance, assume that variances are equal across groups or samples. Tests for homogeneity of variances can be used to verify that assumption and for pooling of data from different sources to yield an improved estimated variance. Tests for homogeneity of variances can be used usually under assumption of normal distributions or nearly normal distributions. In this paper some tests for homogeneity of variances are examined under the null hypothesis and under the alternative, for various sample sizes, for various symmetric and asymmetric distributions. Monte Carlo simulations has been used for this. In this paper the following procedures have been analyzed: Levene’s test, Brown-Forsythe test, Fligner-Killeen test and O’Brien test.

Key words: tests for homogeneity of variances, Monte Carlo, Levene’s test, Brown-Forsythe test, Fligner-Killeen test, O’Brien test.

HELENA GASPARS-WIELOCH

THE DISCRETE LOCATION PROBLEM FOR A CHAIN OF HOMOGENEOUS FACILITIES

1. INTRODUCTION

Location problems arouse interest of many people, but everyone formulates assumptions in a different way. In some papers a continuous approach is presented [2], [11], in others – a discrete one [5]. Some authors are considering the possibility to locate facilities in a linear market, e.g. along the beach [5], [11], others concentrate on bi-dimensional problems, but simultaneously demonstrate that they can be reducible to a one-dimensional case [3]. Objectives are also diverse. The goal of the research may consist, for instance, in maximizing the difference between the demand captured and the demand lost [13], in minimizing the distance between the facility and the highest number of potential clients [11] or in minimizing the total cost both for building and operating facilities as well as for servicing the customer demands [2]. The cases discussed in the literature concern both the location of facilities belonging to one proprietor [4] as well as the location of competitive facilities [1], [5]. In some articles the location is established for a set of similar facilities [5], [10], in others – for a complex of different facilities [14]. In some models location decisions are only considered [5], while other models incorporate additional decisions concerning, e.g. price and production levels [4]. We can find examples where facilities capacity, which may be treated as limited or not, is recognized as a decisive factor influencing the optimal solution [1], [2].

The main objective of this paper is to formulate models determining the best location strategy for a chain of homogeneous facilities, which belong to a single owner. Therefore, the model's goal is to maximize his or her total profit and not the profit of each facility separately. In the research facilities are defined as goods or services distribution centers (e.g. shops, restaurants, petrol stations, post offices, repair services) and are represented as points, not lines (e.g. railways) or polygons. Services are not provided virtually. The facilities' homogeneity means that they offer the same products at the same price, with the same quality of services and with similar opening hours. Although there are lots of contributions devoted to a continuous approach [8], [12], many researchers stress that the analysis of a region as a homogeneous space without geographical barriers is not appropriate [7]. That is why formulated models should be applicable to a discrete approach. It usually means that the location for p facilities is chosen from a set of n possible places (where $p < n$). Here, each potential facility is assigned to a particular location which has already been investigated from an

administrative and transport point of view. Therefore, the decision-maker does not care about the place where to build a facility, but which facilities should be built. In other words, the goal is to choose, from a set of n potential facilities, that subset of p facilities which allows the objective function to achieve the most profitable value. Sections 2-4 describe three possible models designed to establish the best location strategy in view of assumptions given above. Section 5 presents suggestions how to deal with research limitations. Section 6 presents a conclusion.

2. FIRST APPROACH – LOCATION PROBLEM VERSUS RESOURCES ALLOCATION PROBLEM

The comparison of two different issues – the resources allocation problem and the location problem – reveals their significant similarities. In the first case the target is to distribute scarce resources among alternative activities. In the second one, the objective is to assign a limited number of facilities (in view of a given budgetary constraint) to different delimited areas (e.g. quarters, cities, regions). Although this analogy occurs, we can not recognize that the facilities considered in location issues are as totally homogeneous as the resources are in resources allocation problems, because even homogeneous facilities in terms of products, services and prices, differ from each other in their location! That is why, in contrast to resources allocation, in locations problems not only the number of facilities activated at each area has to be defined, but also the combination, which entails a significant growth of the number of possible solutions. Consider n regions and m_j potential facilities in region j . If we want to know how many facilities should be built in n regions we choose one of

$$\prod_{j=1}^n (m_j + 1) \quad (1)$$

possible strategies, but if we also consider the combination, the number of variants rises to

$$2^{\left(\sum_{j=1}^n m_j\right)} \quad (2)$$

Additionally, the location factor influences the number of combinations for which the profit has to be estimated. When regions are far from each other and profits generated by a facility depend only on other facilities built at the same area, the number is equal to

$$\sum_{j=1}^n 2^{m_j}. \quad (3)$$

However, when regions are close to each other, the facility activation may influence profits generated by facilities opened both in the same area as well as in adjacent areas. Then, the profit estimation is required for more cases:

$$2^{\left(\prod_{j=1}^n m_j\right)}.$$

(4)

If in the optimization tasks concerning resources allocation, the total quantity of resources to allocate is not set in advance, the problem may be solved with the aid of the dynamic programming or some simplified procedures such as the local extrema method (for each activity the quantity connected with the highest profit is chosen) or the marginal profits method (resources are allocated as long as marginal profits are non-negative). The last method finds applications only in non-increasing marginal profits functions' cases. Methods mentioned above, except for local extrema method, are also used when the quantity to allocate is known in advance (then, using the marginal profits procedure, facilities built in the first order are those connected with the highest marginal profit). Algorithms characteristic of resources allocation issue may be applied to facilities location on condition that for p_j facilities activated in region j the combination connected with the highest profit is the only one to be taken into account. Moreover, these methods may be used only if interdependency between facilities built in different regions does not occur.

Table 1

Estimated cumulative profits for each region

p_j	Region A		Region B		Region C	
	set of facilities	profit	set of facilities	profit	set of facilities	profit
1	A1	3	B1	3,5	C1	3
	A2	4	–	–	C2	2
	–	–	–	–	C3	3
2	A1, A2	5	–	–	C1, C2	4
	–	–	–	–	C1, C3	6
	–	–	–	–	C2, C3	3,5
3	–	–	–	–	C1, C2, C3	5,5

Source: own calculations.

Table 1 presents cumulative profits estimated for regions A, B, C on the assumption that $m_1 = 2, m_2 = 1, m_3 = 3$. Symbols A1, A2, B1, C1, C2, C3 design potential facilities. Profits depend on the number and set of facilities activated. Cumulative profits are estimated for all possible combinations and the best ones are in bold.

Marginal profits are given in table 2. They are calculated according to the following formula:

$$c'_j(p_j) = c_j(p_j) - c_j(p_j - 1) \quad p_j = 1, 2, \dots, m_j$$

(5)

where $c_j(p_j)$ signifies the estimated profit for the best combination of p_j facilities activated in region j .

Table 2

Marginal profits calculated for the best variants

p_j	Region A		Region B		Region C	
	set of facilities	profit	set of facilities	profit	set of facilities	profit
1	–	–	B1	3,5	C1	3
	A2	4	–	–	–	–
	–	–	–	–	C3	3
2	A1, A2	1	–	–	–	–
	–	–	–	–	C1, C3	3
	–	–	–	–	–	–
3	–	–	–	–	C1, C2, C3	–0,5

Source: own calculations.

If the total number of facilities activated is not set in advance, the optimal strategy consists in choosing all facilities except for C2. If only four facilities may be activated, A2, B1, C1 and C3 should be selected.

The problem can be expressed as

$$\sum_{j=1}^n c_j(p_j) \rightarrow \max \quad (6)$$

$$0 \leq p_j \leq m_j \quad j = 1, 2, \dots, n \quad (7)$$

$$\sum_{j=1}^n p_j = L \leq \sum_{j=1}^n m_j \quad (8)$$

where m_j is defined as the number of potential facilities in region j and L stands for the desired total number of facilities activated. The constraint (8) is optional.

The approach presented in Section 2 is relatively simple, but requires profit estimation for all combinations.

3. SECOND APPROACH – PROFIT FUNCTION FOR EACH FACILITY

In connection with the necessity of estimating the profit for plenty of variants in the first approach, the author is considering the possibility of deriving profit functions for each potential facility. Then, the maximization of their sum may constitute the

objective function. The profit generated by a given facility certainly depends on its territory served called also range of coverage or influence area. Factors determining this territory can be divided into two groups. The first one consists of factors known or relatively easy to predict before the decision making (e.g. place attractiveness, population density and existing competition) and allows us to establish the original range of coverage (W_{ij}^o). This area may be reduced after taking into account factors belonging to the second group and relating to the activation of other potential facilities from the same chain as well as to the distances between them and a particular facility. These data are known just when the optimal solution is found. The final range of coverage may be for instance defined as

$$w_{ij} = \max \left\{ 0, W_{ij}^o x_{ij} - \left(-x_{ij} + \exp \left(\sum_{l=1}^n \sum_{k=1}^{m_l} s_{ij}^{kl} x_{kl} \right) \right) \right\} \quad (9)$$

where x_{ij} is equal to 1 if the facility i is activated in region j and 0 otherwise, s_{ij}^{kl} is the so-called nonnegative shrinking coefficient representing the impact the activation of facility kl has on the territory served by facility ij . When facilities are far from each other the shrinking coefficient is equal to zero. The range of coverage is not treated as a circle whose radius decreases when a new facility is activated in the vicinity, because the territory served shrinks only near the new facility and not from all directions. Depending on the combination of facilities selected to be built and their location, the influence area will be equal to one of the following values:

$$(x_{ij} = 0) \Rightarrow (w_{ij} = 0) \quad (10)$$

$$\left((x_{ij} = 1) \wedge (\forall x_{kl} \neq x_{ij} : x_{kl} = 0) \right) \Rightarrow (w_{ij} = W_{ij}^o) \quad (11)$$

$$\left((x_{ij} = 1) \wedge (\exists x_{kl} \neq x_{ij} : x_{kl} = 1 \wedge s_{ij}^{kl} = 0) \right) \Rightarrow (w_{ij} = W_{ij}^o) \quad (12)$$

$$\left((x_{ij} = 1) \wedge (\exists x_{kl} \neq x_{ij} : x_{kl} = 1 \wedge s_{ij}^{kl} > 0) \right) \Rightarrow (w_{ij} < W_{ij}^o) \quad (13)$$

The equation (9) is correct only if possible intersections of original ranges of coverage occur between two facilities (Figure 1).

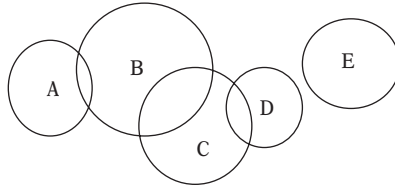


Figure 1. Original ranges of coverage for five potential facilities

Otherwise, the impact of facilities activation could be overestimated. Therefore, the formula would have to be modified by adding shrinking coefficients taking into account the simultaneous effect of the activation of more facilities.

Given the territory served by facility ij we can define its profit (P_{ij}) as a function of w_{ij} . The second model can emerge in the form of the following equations:

$$\sum_{j=1}^n \sum_{i=1}^{m_j} P_{ij} x_{ij} \rightarrow \max \quad (14)$$

$$P_{ij} = f(w_{ij}) \quad i = 1, 2, \dots, m_j, \quad j = 1, 2, \dots, n \quad (15)$$

$$w_{ij} = \max \left\{ 0, W_{ij}^o x_{ij} - \left(-x_{ij} + \exp \left(\sum_{l=1}^n \sum_{k=1}^{m_l} s_{ij}^{kl} x_{kl} \right) \right) \right\} \quad (16)$$

$$i = 1, 2, \dots, m_j, \quad j = 1, 2, \dots, n$$

$$\sum_{j=1}^n \sum_{i=1}^{m_j} x_{ij} = L \leq \sum_{j=1}^n m_j \quad (17)$$

$$x_{ij}, x_{kl} \in \{0, 1\} \quad (18)$$

Since people living far from a facility will not necessarily patronize this one, it is possible to recognize that along with the growth of the range of coverage profits P_{ij} generated by a particular facility increase as well but at a decreasing speed. Thus, the author recommends applying such analytical forms for the equation (15) as hyperbola: $P_{ij} = \alpha + \beta/w_{ij}$ with a constant positive and a coefficient negative, Törnquist function: $P_{ij} = (\alpha \cdot w_{ij})/(\beta + w_{ij})$ with parameters positive, or power function: $P_{ij} = \alpha \cdot w_{ij}^\beta$, where α is positive and β belongs to the interval (0,1).

4. THIRD APPROACH – PROFIT FUNCTION FOR EACH MARKET AREA

The second approach does not require estimating final profits for numerous combinations, but establishing a suitable analytical form and parameters for each territory served and profit function, which may cause difficulties. Therefore, one more algorithm with the following steps is suggested:

1. Specify the area with potential facilities.
2. Split this area into K smaller market areas (Figure 2).

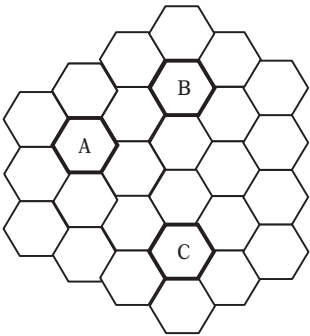


Figure 2. Splitted area with potential facilities

3. Determine for each potential facility the original range of coverage as a set of market areas (Figure 3).

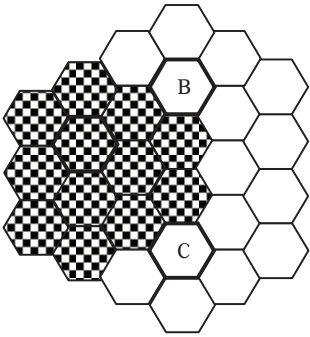


Figure 3. Facility A: the original range of coverage

4. Estimate profits generated by potential facilities from each market area assuming that only one facility will be activated.

5. Distinguish those market areas which belong to more than one original territory served (Figure 4).

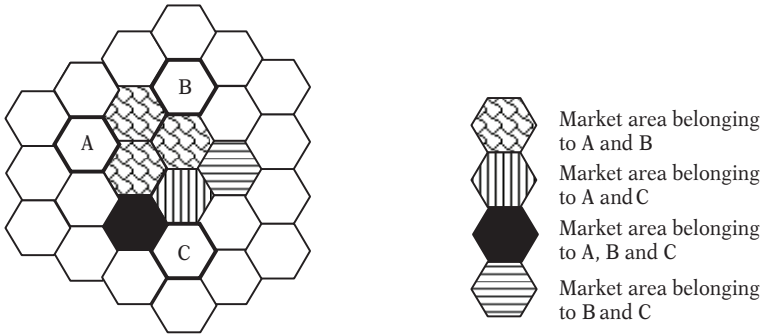


Figure 4. Market areas belonging to more than one original territory served

6. Estimate shares in profits when more than one facility takes advantages of a particular market area.

7. Using a suitable optimization model find the set of L facilities maximizing the total profit generated from all market areas.

It is recommended to present the split area as a hexagonal honeycomb proposed by A. Lösch [6], because hexagons cover the area without leaving holes. Steps 3 and 5 allow us to obtain a diagram similar to a Voronoi diagram, also called Dirichlet tessellation [15]. Using Euclidean distance as the only criterion to delimit territories served for each facility is impossible, because customers choose the store considering a trade-off between proximity and store attractiveness [7]. That is why the diagram does not resemble an ordinary Voronoi diagram, but a multiplicatively weighted Voronoi diagram. The third optimization model is presented below:

$$\sum_{k=1}^K h_k \rightarrow \max \quad (19)$$

$$h_k = \begin{cases} 0 & \text{if } \forall x_{ij}^k : x_{ij}^k = 0 \\ \sum_{j=1}^n \sum_{i=1}^{m_j} h_{ij}^k s_{ij}^k x_{ij}^k & \text{if } \exists x_{ij}^k : x_{ij}^k = 1 \end{cases} \quad k = 1, 2, \dots, K \quad (20)$$

$$s_{ij}^k = f(x_{11}^k, x_{21}^k, \dots, x_{m_1 1}^k, x_{12}^k, x_{22}^k, \dots, x_{m_2 2}^k, \dots, x_{m_n n}^k) \quad (21)$$

$$i = 1, 2, \dots, m_j, j = 1, 2, \dots, n, k = 1, 2, \dots, K$$

$$x_{ij} = \begin{cases} 0 & \text{if } \sum_{k=1}^K x_{ij}^k = 0 \\ 1 & \text{if } \sum_{k=1}^K x_{ij}^k > 0 \end{cases} \quad i = 1, 2, \dots, m_j, j = 1, 2, \dots, n \quad (22)$$

$$\sum_{j=1}^n \sum_{i=1}^{m_j} x_{ij} = L \leq \sum_{j=1}^n m_j \quad (23)$$

$$s_{ij}^k \in [0, 1] \quad (24)$$

$$x_{ij}, x_{ij}^k \in \{0, 1\} \quad (25)$$

where h_k means the estimated profit generated by the market area k . The variable x_{ij}^k is equal to 1 when the facility ij , whose original territory served covers this market area, is activated. The estimated profit generated from the market area k by the facility ij being the only one taking advantages of this market area is given by h_{ij}^k and s_{ij}^k , signifies facility share in profits from the market area. The variable x_{ij} is equal to 1 if the facility i is activated in region j and 0 otherwise.

Notice that shares in profits depend on the set of facilities activated. Assume that the market area $k = 1$ belongs to the original range of coverage of two potential facilities: $(i = 1, j = 1)$ and $(i = 2, j = 1)$. If only one of them is built, the estimated profit is equal to $h_{11}^1 = 7$ or $h_{21}^1 = 6$. If the two facilities are activated, shares will amount to $s_{11}^1 = 3/4$ and $s_{21}^1 = 1/2$. For the case presented above, the equation (20) comes in the form:

$$h_1 = \begin{cases} 0 & \text{if } x_{11}^1 = 0, x_{21}^1 = 0 \\ 7 \cdot 1 \cdot x_{11} & \text{if } x_{11}^1 = 1, x_{21}^1 = 0 \\ 6 \cdot 1 \cdot x_{21} & \text{if } x_{11}^1 = 0, x_{21}^1 = 1 \\ 7 \cdot (3/4) \cdot x_{11} + 6 \cdot (1/2) \cdot x_{21} & \text{if } x_{11}^1 = 1, x_{21}^1 = 1 \end{cases} \quad (26)$$

The sum of shares for a given market area does not need to be equal to 1, because shares signify a part of profits generated by a particular facility and not a part of the profit achieved together. When the area is served by more than one facility, the global profit from this area is usually higher than the profit generated from it by the only facility activated.

5. PARAMETERS ESTIMATION

The application of models presented above is quite simple apart from parameters estimation which concerns the territory served as well as profits and shares in profits. There are parameters that should describe the influence of exogenous factors relating to the population, place attractiveness, existing competition and accessible means of transport (e.g. W_{ij}^o, h_{ij}^k) and parameters showing how a given facility may depend on other facilities from the same chain (e.g. s_{ij}^{kl}, s_{ij}^k). Notice that depending on how travel time or distances are measured parameters amount to different values. Ponsard [9] suggests applying an Euclidean, rectangular, peripheral distance or a network, but others recommend especially the last measure saying that human mobility is usually limited by the transport network. In addition, a network space, in contrast to Euclidean or rectangular distances, reduces the problem to one dimension [2]. Models presented by Sadahiro [10], Dasci and Laporte [4] may be also very useful. Sadahiro states that the accessibility and concentration measures may be established on the basis of opening hours and locations of each facility, which may help to estimate shares in profits. Dasci and Laporte show how to find the market boundary for two neighbor stores under the assumption that customers patronize that one which provides the minimum delivered price (products' price + cost of transport). Many other approaches relating to patronage probability calculation are presented in literature (e.g. [1], [7]).

6. CONCLUSION

This paper presents three possible approaches concerning the optimal location of homogeneous facilities belonging to one proprietor. In the first approach it is

demonstrated that methods used in resources allocation problem may also be applied to location issues but only if interactions occur between facilities built in the same region. The second one emphasizes the range of coverage as an important factor determining profits generated by facilities, and in the third approach Lösch's theory and Voronoi diagrams are used to maximize the total profit from market areas. Depending on available data, one of these models can be applied. They allow taking into consideration interactions between facilities and they are quite simple. Moreover, models are flexible because their equations may have other analytical forms than those presented in the paper. In the further investigation that could be carried out, parameters may be replaced with random variables with known distribution. A natural extension is to enrich the research by a temporal facet where cumulative values of parameters are estimated for different periods. Such an approach allows to consider the evolution of phenomena influencing profits generated by a facility as well as taking into account new significant factors that may appear in the future (e.g. new competition). Such an approach also allows to establish an optimal location strategy for a given moment.

Uniwersytet Ekonomiczny w Poznaniu

REFERENCES

- [1] Brandeau M.L., Chiu S.S., [1994], *Location of competing facilities in useroptimizing environment with market externalities*, TS, 28 (2), 125-140.
- [2] Brimberg J., Korach E., Eden-Chaim M., Mehrez A., [2001], *The capacited p-facility location problem on the real line*, „International Transactions of Operational Research”, 8, 727-738.
- [3] Brimberg J., Love R.F., [1998], *Solving a class of two-dimensional uncapacited location-allocation problems by dynamic programming*, „Operations Research”, 46, 702-709.
- [4] Dasci A., Laporte G., [2004], *Location and pricing decisions of a multistore monopoly in a spatial market*, „Journal of Regional Science”, 44, 489-515.
- [5] Hotelling H., [1929], *Stability in competition*, „Economic Journal”, 39, 41-57.
- [6] Lösch A., [1940], *Die räumliche Ordnung der Wirtschaft*, Jena.
- [7] Mendes A.B., Themido I.H., [2004], *Multi-outlet retail site location assessment*, „International Transactions in Operational Research”, 11, 1-18.
- [8] Palander T., [1935], *Beiträge zur Standorttheorie*, Almqvist & Wiksell, Uppsala.
- [9] Ponsard C., [1979], *Economie urbaine et espaces métriques*, „Sistemi Urbani”, 1, 123-135.
- [10] Sadahiro Y., [2005], *Spatiotemporal analysis of the distribution of urban facilities in terms of accessibility*, „Papers in Regional Science”, 84 (1), 61-84.
- [11] Samuelson W.F., Marks S.G., [2002], *Managerial Economics*, Wiley Publishing CO, New York.
- [12] Steinhaus H., [1956], *Kalejdoskop Matematyczny*, PWN, Warszawa.
- [13] Stevens B.H., [1961], *An application of game theory to a problem in location strategy*, „Papers and Proceedings in Regional Science Association”, 7.
- [14] Thünen J.H., [1826], *Der isolierte Staat In Beziehung auf Landwirtschaft und Nationalökonomie*, Hamburg.
- [15] Voronoi G.F., [1908], *Nouvelles applications des paramètres continus a la théorie des formes quadratiques. Deuxième mémoire. Recherche sur les paralléloèdres primitifs*, „Journal für Reine und Angewandte Mathematik”, 134, 198-287.

DYSKRETNE ZAGADNIENIE LOKALIZACYJNE DLA SIECI JEDNORODNYCH OBIEKTÓW

Streszczenie

Autorka pracy wymienia możliwe kryteria podziału zagadnień lokalizacyjnych opisanych w literaturze, a następnie przedstawia propozycję trzech dyskretnych modeli optymalizacyjnych, które mogą znaleźć zastosowanie przy opracowywaniu projektu uruchomienia sieci jednorodnych obiektów należących do jednego właściciela. Celem zadań jest maksymalizacja całkowitego zysku osiągniętego przez tegoż właściciela, a nie maksymalizacja zysku zrealizowanego przez każdy obiekt osobno. Autorka ukazuje powiązanie pierwszego modelu z zagadnieniem alokacji zasobów. Wpływ odległości pomiędzy obiektami na obszar zasięgu poszczególnych obiektów został w szczególności uwzględniony w drugim i trzecim modelu optymalizacyjnym. W ostatnim zadaniu wykorzystano założenia Lösch'a i Voronoi.

Słowa kluczowe: Zagadnienie lokalizacyjne, Dyskretny model optymalizacyjny, Maksymalizacja zysku, Jednorodne obiekty, Obszar zasięgu, Diagram Voronoja (tessellacja Dirichleta), Zagadnienie rozdziału zasobu, Metoda ekstremów lokalnych, Metoda zysków krańcowych.

THE DISCRETE LOCATION PROBLEM FOR A CHAIN OF HOMOGENEOUS FACILITIES

Summary

The beginning of the article is devoted to a review of different location problems discussed in the literature. In the main part of this contribution the author presents and compares three discrete optimization models that may be useful for decision-makers considering the construction and activation of a chain of homogeneous facilities belonging to one proprietor. The models goal is to maximize his or her total profit and not the gain of each facility separately. The author shows the connection of the first model with the resources allocation problem. The influence of the distance between facilities on their territory served is emphasized particularly in the second and third approach. The last model is partially based on Lösch's and Voronoi's principles.

Key words: Location problem, Discrete optimization models, Profit maximization, Homogeneous facilities, Territory served (range of coverage, influence area), Voronoi diagram (Dirichlet tessellation), Resources allocation problem, Local extrema method, Marginal profits method.

ANNA ZAMOJSKA

ZASTOSOWANIE METODY DEA W KLASYFIKACJI FUNDUSZY INWESTYCYJNYCH

1. WSTĘP

Analiza i ocena wyników osiągniętych przez fundusze inwestycyjne jest jednym z podstawowych etapów w procesie oceny efektywności zarządzania portfelem inwestycyjnym funduszu. Jedno z najstarszych, obecnie stosowanych podejść do oceny efektywności oparte jest na następujących wskaźnikach: Jensena [4], Sharpe'a [12] i Treynora [13]. Wskaźniki te liczone są w oparciu o wyniki estymacji Modelu Wyceny Aktywów Kapitałowych (CAPM). Nie zawsze istnieje możliwość ich pełnej interpretacji, co wynika z faktu, iż podstawową wadą wymienionych wskaźników jest trudność spełnienia bardzo rygorystycznych założeń modelu CAPM odnośnie struktury rynku walorów oraz postaci rozkładu ich stóp zwrotu. Najstarsze, klasyczne podejście do mierzenia efektywności zarządzania portfelem inwestycji funduszu, wykorzystuje zazwyczaj takie elementy, jak stopa zwrotu i ryzyko. Doświadczenia ostatnich lat pokazują jednak, że te dwa parametry są niewystarczające, jeśli chcemy porównać między sobą wyniki osiągane przez poszczególne fundusze funkcjonujące na rynku. Jednym z istotnych elementów, który należy uwzględnić to koszty, jakie ponoszone są w związku z uczestnictwem w danym funduszu i tym samym mają wpływ na poziom zysków osiągniętych przez inwestorów. Nawet niewielka różnica w zysku netto osiąganym w krótkim okresie czasu, prowadzi do bardzo dużych różnic w łącznym zysku netto w długim okresie czasu. W 1997 r. Murthi i in. zaproponowali metodę Data Envelopment Analysis w skrócie DEA¹ do oceny wyników osiągniętych przez fundusze [5]. Za pomocą metody DEA określany jest poziom efektywności funduszu względem innych funduszy z badanej próby.

Artykuł prezentuje zastosowanie metody DEA do mierzenia efektywności zarządzania portfelem inwestycyjnym funduszy akcyjnych funkcjonujących na polskim rynku kapitałowym. W pierwszej części artykułu zawarto teoretyczne podstawy metody DEA, następnie przedstawiono postacie prymalne i dualne zagadnień programowania liniowego. W części drugiej opisano stosowane warianty metody DEA, które zostały wykorzystane do pomiaru efektywności funduszy, uwzględniające odpowiednio ryzyko całkowite i systematyczne, stopy zwrotu oraz koszty uczestnictwa w funduszy, z uwzględnieniem krótkiego i długiego horyzontu inwestowania. W trzeciej części artykułu zamieszczono otrzymane wartości poziomu efektywności badanych funduszy akcyjnych oraz wyniki

¹ Polskie tłumaczenia Data Envelopment Analysis (DEA) to: Metoda Obwiedni Danych [6], Metoda Otoczki Danych [8], Metoda Analiz Granicznych [7].

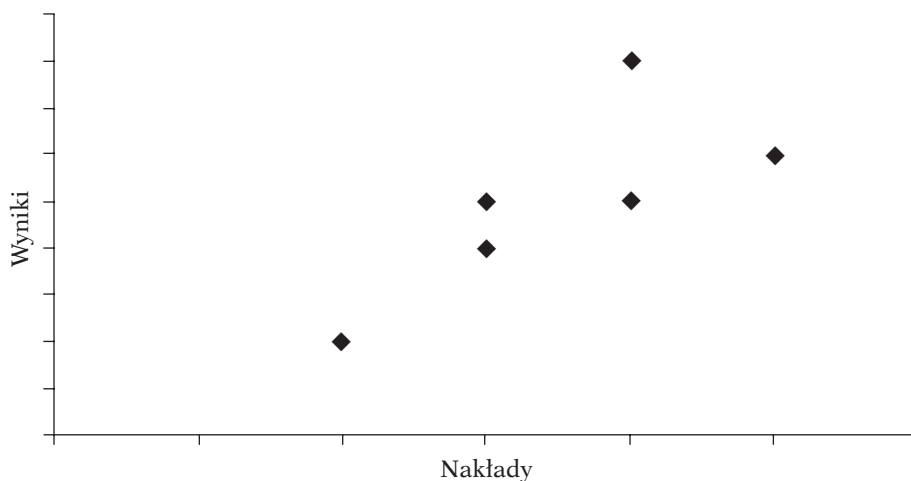
przeprowadzonej analizy porównawczej między metodą DEA a klasycznymi wskaźnikami Jensena, Treynora i Sharpe'a.

2. TEORETYCZNE PODSTAWY METODY DEA

DEA jest jedną z technik badań operacyjnych stosowaną w analizie problemu efektywności [2] i należy do grupy metod nieparametrycznych. Metoda ta jest dość powszechnie stosowana na świecie w różnych obszarach gospodarki. Początkowo stosowana była do oceny efektywności działalności podmiotów sektora publicznego i jednostek nieprzynoszących zysków, takich jak szkoły czy placówki medyczne [1]. W kolejnych latach metodę DEA zaczęto stosować do oceny efektywności banków [2], [14], przedsiębiorstw branży energetycznej [2], funduszy inwestycyjnych [1], [3], [5], [15]. W Polsce metoda ta jest także stosowana i istnieje wiele publikacji prezentujących wyniki przeprowadzonych badań w różnych obszarach, takich jak: działalność banków [8], sektor energetyczny [8] czy biblioteki [6]. W przypadku funduszy inwestycyjnych metodę DEA w ocenie efektywności funduszy jako pierwszy zastosował Murthi w 1997 r. [5], kolejno zaaplikowana została przez Basso i Funari [1]. W obu pracach jako wyniki wykorzystano wyniki funduszy, a jako nakłady ryzyko całkowite mierzone odchyleniem standardowym, ryzyko systematyczne mierzone parametrem beta oraz koszty transakcyjne. Wykorzystanie wartości narażonej na ryzyko (VaR) pozwala na uwzględnienie w nakładach w modelu DEA pojawiających się często w szeregach stóp zwrotu funduszy wartości ekstremalnych [9].

Podstawą metody DEA jest współczynnik produktywności Debreu-Farella wyrażony jako stosunek jednego nakładu i jednego wyniku [1], który został uogólniony na przypadek wielowymiarowy, czyli wielu wyników i wielu nakładów. Przedmiotem analizy w metodzie DEA jest określenie poziomu efektywności, z jaką podmiot podejmujący decyzję DMU (z ang. *decision making units*) transformuje posiadane nakłady na wyniki. Za pomocą metody DEA wyznaczana jest granica efektywności zbioru możliwości produkcyjnych. Obiekty znajdujące się na tej granicy przyjmują wartość współczynnika efektywności równą 1, natomiast wartość tego współczynnika w przypadku obiektów leżących poniżej tej granicy jest mniejsza od jedności. Współczynnik przyjmuje wartości z przedziału $(0,1)$, a różnica wartości względem 1 określa rozmiar nieefektywności pojedynczego obiektu, ponieważ metoda DEA pozwala określić, jaki jest poziom efektywności wybranego obiektu względem pozostałych obiektów w analizowanej próbie.

Podstawową zaletą stosowania metody DEA jest to, że nie wymaga ona formułowania postaci funkcyjnej zależności między nakładami a wynikami.



Wykres 1. Zbiór możliwości produkcyjnych

Źródło: opracowanie własne.

Przyjmijmy następujące oznaczenia:

$j = 1, 2, \dots, n$ obiekty,

$r = 1, 2, \dots, t$ wyniki,

$i = 1, 2, \dots, m$ nakłady,

y_{rj} wielkość wyniku r na jednostkę j ,

x_{ij} wielkość nakładu i na jednostkę j ,

u_r waga wyniku r ,

v_i waga nakładu i .

Miara efektywności DEA dla jednostki (obektu) analizy j jest obliczana jako stosunek sumy ważonych wyników do sumy ważonych nakładów postaci [1]:

$$h_j = \frac{\sum_{r=1}^t u_r y_{rj}}{\sum_{i=1}^m v_i x_{ij}}. \quad (1)$$

W przypadku obiektów efektywnych wartość współczynnika h jest równa 1 (obiekt znajduje się na granicy zbioru możliwości produkcyjnych (inwestycyjnych, granica efektywna), natomiast w przypadków obiektów nieefektywnych jego wartość jest mniejsza od 1 i obiekt znajduje się poniżej granicy efektywności. Wielkość wag określana jest w procesie optymalizacji wartości współczynnika h , dla każdej jednostki analizy osobno. Ilorazowe zagadnienie programowania matematycznego dla wybranego obiektu j_0 ma postać:

$$\max_{\{v_i, u_r\}} h_{j0} = \frac{\sum_{r=1}^t u_r y_{rj0}}{\sum_{i=1}^m v_i x_{ij0}} \quad (2)$$

p.w.

$$\begin{aligned} \frac{\sum_{r=1}^t u_r y_{rj}}{\sum_{i=1}^m v_i x_{ij}} &\leq 1, \quad j = 1, \dots, n, \\ u_r &\geq \varepsilon, \quad r = 1, \dots, t, \\ v_i &\geq \varepsilon, \quad i = 1, \dots, m, \end{aligned} \quad (3)$$

Ostatnie dwa warunki oznaczają, że wagi nakładów i wyników są zawsze dodatnie i większe od zera. Optymalna wartość funkcji celu (2) jest miarą efektywności dla danego obiektu analizy. Proces wyznaczania wartości współczynnika efektywności h , poprzez optymalizację modelu postaci (2) i (3) powtarzany jest dla każdego obiektu należącego do analizowanej próby. W procesie wyznaczania optymalnej wartości funkcji celu wyznaczane są wartości wag u_r i v_i , co pozwala agregować nakłady i wyniki bez uwzględniania struktury preferencji osoby podejmującej decyzję. Ponieważ wartość optymalna funkcji celu zmienia się dla kolejnych obiektów, to obliczone wartości wag wyników i nakładów odzwierciedlają ich specyfikę.

Zagadnienie ilorazowego programowania matematycznego określone przez (2) i (3) można przekształcić w alternatywny problem programowania liniowego, nakładając ograniczenie postaci $\sum v_i x_{ij0} = 1$. Jest to liniowy model CCR (od nazwisk autorów Charnes, Cooper i Rhodes) nazywany modelem zorientowanym na nakłady (z ang. *input-oriented linear model*) [1].

Funkcja celu modelu CCR dla wybranego obiektu $j0$ ma postać:

$$\max \sum_{r=1}^t u_r y_{rj0} \quad (4)$$

p.w.

$$\begin{aligned} \sum_{i=1}^m v_i x_{ij0} &= 1, \\ \sum_{r=1}^t u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} &\leq 0, \quad j = 1, \dots, n, \\ -u_r &\leq -\varepsilon, \quad r = 1, \dots, t, \\ -v_i &\leq -\varepsilon, \quad i = 1, \dots, m. \end{aligned} \quad (5)$$

Ten problem ma $(t + m)$ zmiennych decyzyjnych, którymi są wagi u i v , wybierane tak, aby maksymalizować efektywność wybranego podmiotu j_0 . Liczba warunków ograniczających jest równa $(n + t + m + 1)$.

Zagadnienie dualne do postaci pierwotnej (4) i (5) ma postać:

$$\min_{\{\theta, \lambda, s_x, s_y\}} \theta - \varepsilon(I^T s_x + I^T s_y) \quad (6)$$

p.w.

$$\begin{aligned} X\lambda - \theta x_0 + s_x &= 0 \\ Y\lambda - y_0 - s_y &= 0 \\ \lambda &\geq 0, s_x \geq 0, s_y \geq 0 \\ \theta \in R, \lambda \in R^n, s_x \in R^m, s_y \in R^p. \end{aligned} \quad (7)$$

Na zmienne modelu θ i λ_j nie są nałożone żadne ograniczenia, co do ich wartości. Zmienna θ jest interpretowana jako wielkość wszystkich nakładów przypadająca na jednostkę produkcji nieefektywnego obiektu j_0 , niezbędna do osiągnięcia poziomu efektywności równej jedności [3]. Zbiór zmiennych λ_j zdefiniowany jest dla każdej jednostki analizy i definiuje punkt przestrzeni obwiedni (wartości granicznych) utworzony jako wypukła kombinacja liniowa obiektów efektywnych.

3. ZASTOSOWANIE METODY DEA DO OCENY EFEKTYWNOŚCI FUNDUSZY

Ocena efektywności funduszy inwestycyjnych jest przedmiotem wielu badań teoretycznych i empirycznych. W klasycznym, parametrycznym podejściu do oceny ich efektywności wymagane jest spełnienie szeregu założeń, zwykle trudnych do spełnienia. Obecnie można wskazać, co najmniej 100 wskaźników oceny efektywności funduszy, z których żaden w sposób kompleksowy nie ocenia efektywności funduszu, tak by móc jednoznacznie ocenić poziom efektywności funduszu. Metoda DEA, która jak wcześniej wspomniano została zaadaptowana do oceny wyników funduszy inwestycyjnych [5] traktuje każdy pojedynczy fundusz jako DMU [9]. W porównaniu z klasycznym podejściem zastosowanie metody DEA w ocenie efektywności funduszy inwestycyjnych ma następujące zalety: jest to podejście nieparametryczne, nie są wymagane żadne założenia odnośnie wielkości benchmarku, można uwzględnić jednocześnie kilka różnych zestawów zmiennych odnośnie wielkości nakładów i wyników, łatwa w użyciu, niezwykle prosta w interpretacji, dająca jednoznaczny wynik ostateczny w postaci rankingu funduszy.

Do oceny efektywności funduszy inwestycyjnych wykorzystać można model zorientowany na nakłady (CCR). W proponowanym modelu² wynikami są oczekiwane stopy zwrotu lub nadwyżki oczekiwanych stóp zwrotu oraz współczynnik asymetrii stóp

² Model został wykorzystany do oceny efektywności rynku funduszy we Włoszech [1].

zwrotu, natomiast nakładami są ryzyko oraz koszty, jakie ponosi inwestor (koszty ponoszone przy nabywaniu jednostek funduszu oraz ich wykupywaniu lub umarzaniu). Interpretacja tak zdefiniowanych wyników i nakładów ma swoje uzasadnienie w specyfice inwestycji finansowej, jaką są fundusze inwestycyjne. Celem ich działalności jest osiągnięcie jak najwyższej stopy zwrotu ze środków powierzonych im przez indywidualnych inwestorów, i tym samym jest to element stanowiący efekt końcowy działań podejmowanych przez fundusz. Ponieważ decyzja o wyborze funduszu nie jest podejmowana w oparciu o tylko jedną stopę zwrotu, ale najczęściej w oparciu o szereg czasowy stóp zwrotu uzasadnione jest wprowadzenie współczynnika asymetrii stóp zwrotu, który jeśli jest dodatni świadczy o dobrym zarządzaniu portfelem inwestycyjnym funduszu, ponieważ występuje asymetria między ilością stóp zwrotu powyżej i poniżej średniej, ze wskazaniem na te, powyżej, co jest korzystne z punktu widzenia inwestora. Zgodnie z podstawową zasadą obowiązującą na rynku finansowym, osiągnięcie określonego poziomu stopy zwrotu narażone jest na ryzyko, i co istotne im wyższa stopa zwrotu tym większe ryzyko. W konsekwencji oznacza to, że także fundusze obciążone są ryzykiem zarządzania portfelem, co należy zinterpretować jako nakład w przypadku stosowania metody DEA. Konstrukcja współczynnika h w oparciu o odpowiednio dobrane wyniki i nakłady sprowadza się do wyznaczenia klasycznych wskaźników oceny efektywności portfela inwestycji, takich jak: współczynnik Treynora (jeden wynik to nadwyżka stopy zwrotu, a jeden nakład to miara ryzyka, jaką jest parametr β) czy Sharpe'a (jeden wynik to nadwyżka stopy zwrotu, a jeden nakład to miara ryzyka, jaką jest odchylenie standardowe stóp zwrotu).

Niech będzie dany zbiór n funduszy i niech wynikiem y_j ($j = 1, 2, \dots, n$) będzie oczekiwana stopa zwrotu $E(R_j)$ lub nadwyżka oczekiwanej stopy zwrotu $[E(R_j) - R_f]$. Należy zauważyć, że pierwszy wybór pozwala zredukować zjawisko ujemnych wartości w wynikach, podczas gdy w drugim obliczane są tradycyjne wskaźniki Treynora, Sharpe'a, Jensena. W pracy zostaną przeprowadzone optymalizacje pierwszego podejścia, czyli oczekiwana stopa zwrotu w trzech podstawowych postaciach modelu, i dodatkowo dla każdej z trzech postaci uwzględnione zostaną po dwa warianty, różniące się między sobą miarą ryzyka.

Model DEA₁₁

W pierwszym modelu DEA₁₁ uwzględniony zostanie jeden wynik (stopa zwrotu funduszu) oraz alternatywnie jedna z miar ryzyka jako nakład (odchylenie standardowe stóp zwrotu, jako miara ryzyka całkowitego DEA_{11c} lub parametr β jako miara ryzyka systematycznego DEA_{11b}).

Model w postaci ilorazowego zagadnienia programowania matematycznego:

$$\max_{\{u, v_i\}} h_0 = \frac{u y_{j0}}{\sum_{i=1}^h v_i x_{ij0}} \quad (8)$$

p.w.

$$\begin{aligned}
\frac{uy_j}{\sum_{i=1}^h v_i x_{ij}} &\leq 1, \quad j = 1, \dots, n, \\
u &\geq \varepsilon, \quad r = 1, \dots, t, \\
v_i &\geq \varepsilon, \quad i = 1, \dots, h.
\end{aligned} \tag{9}$$

Po sprowadzeniu do liniowego zagadnienia model (8) i (9) ma postać:

$$\max \quad uy_{j0} \tag{10}$$

p.w.

$$\begin{aligned}
\sum_{i=1}^h v_i x_{ij0} &= 1, \quad j = 1, \dots, n, \\
uy_j - \sum_{i=1}^h v_i x_{ij} &\leq 0 \\
u &\geq \varepsilon, \quad r = 1, \dots, t, \\
v_i &\geq \varepsilon, \quad i = 1, \dots, h.
\end{aligned} \tag{11}$$

Model DEA₂₁

W drugim modelu DEA₂₁ uwzględniony jest nadal jeden wynik (stopa zwrotu funduszu), natomiast nakłady będą dwa. Nakładami, analogicznie jak w pierwszym modelu, alternatywnie odchylenie standardowe stóp zwrotu lub parametr $\beta(x_{ij})$ i koszty (c_{ij}), jakie ponosi inwestor w związku z nabywaniem i posiadaniem jednostek danego funduszu.

Model w postaci ilorazowego zagadnienia programowania matematycznego:

$$\max_{\{u, v_i, w_i\}} h_0 = \frac{uy_{j0}}{\sum_{i=1}^h v_i x_{ij0} + \sum_{i=1}^k w_i c_{ij0}} \tag{12}$$

p.w.

$$\begin{aligned}
\frac{uy_j}{\sum_{i=1}^h v_i x_{ij} + \sum_{i=1}^k w_i c_{ij}} &\leq 1, \quad j = 1, \dots, n, \\
u &\geq \varepsilon, \quad r = 1, \dots, t, \\
v_i &\geq \varepsilon, \quad i = 1, \dots, h, \\
w_i &\geq \varepsilon, \quad i = 1, \dots, k.
\end{aligned} \tag{13}$$

Po sprowadzeniu do liniowego zagadnienia model (12) i (13) ma postać:

$$\max \quad uy_{j0} \quad (14)$$

p.w.

$$\begin{aligned} \sum_{i=1}^h v_i x_{ij0} + \sum_{i=1}^k w_i c_{ij0} &= 1, \quad j = 1, \dots, n, \\ uy_j - \sum_{i=1}^h v_i x_{ij} - \sum_{i=1}^k w_i c_{ij} &\leq 0 \\ u &\geq \varepsilon, \quad r = 1, \dots, t, \\ v_i &\geq \varepsilon, \quad i = 1, \dots, h, \\ w_i &\geq \varepsilon, \quad i = 1, \dots, k. \end{aligned} \quad (15)$$

Model DEA₂₂

W trzecim modelu DEA₂₂ uwzględniono dwa nakłady analogicznie jak we wcześniejszych modelach oraz dwa wyniki, którymi są stopa zwrotu funduszu (y_j) oraz współczynnik asymetrii stóp zwrotu (d_j).

Model w postaci ilorazowego zagadnienia programowania matematycznego:

$$\max_{\{u_r, v_i, w_i\}} h_0 = \frac{u_1 y_{j0} + u_2 d_{j0}}{\sum_{i=1}^h v_i x_{ij0} + \sum_{i=1}^k w_i c_{ij0}} \quad (16)$$

p.w.

$$\begin{aligned} \frac{u_1 y_j + u_2 d_j}{\sum_{i=1}^h v_i x_{ij} + \sum_{i=1}^k w_i c_{ij}} &\leq 1, \quad j = 1, \dots, n, \\ u &\geq \varepsilon, \quad r = 1, \dots, t, \\ v_i &\geq \varepsilon, \quad i = 1, \dots, h, \\ w_i &\geq \varepsilon, \quad i = 1, \dots, k. \end{aligned} \quad (17)$$

Po sprowadzeniu do liniowego zagadnienia (16) i (17) ma postać:

$$\max (u_1 y_{j0} + u_2 d_{j0}) \quad (18)$$

p.w.

$$\begin{aligned}
 \sum_{i=1}^h v_i x_{ij0} + \sum_{i=1}^k w_i c_{ij0} &= 1, \quad j = 1, \dots, n, \\
 u_1 y_j + u_2 d_j - \sum_{i=1}^h v_i x_{ij} - \sum_{i=1}^k w_i c_{ij} &\leq 0 \\
 u_r &\geq \varepsilon, \quad r = 1, \dots, t, \\
 v_i &\geq \varepsilon, \quad i = 1, \dots, h, \\
 w_i &\geq \varepsilon, \quad i = 1, \dots, k.
 \end{aligned} \tag{19}$$

4. WYNIKI PRZEPROWADZONYCH BADAŃ

Do oceny efektywności funduszy inwestycyjnych wykorzystane zostaną modele zdefiniowane jako DEA₁₁, DEA₂₁, DEA₂₂. W modelu³ jako wyniki wykorzystane zostaną średnie stopy zwrotu funduszy inwestycyjnych, natomiast jako nakłady – klasyczne miary ryzyka, takie jak odchylenie standardowe stóp zwrotu lub parametr β a także koszty, jakie ponosi inwestor w związku z uczestnictwem w danym funduszu inwestycyjnym (koszty ponoszone przy nabywaniu jednostek funduszu, ich wykupywaniu lub umarzaniu).

Przedmiotem badania były akcyjne fundusze inwestycyjne funkcjonujące na polskim rynku kapitałowym. Do badania wybrane zostały dwa okresy, które odpowiednio odzwierciedlają krótki (roczny) i długi horyzont inwestycyjny. Badanie zostało przeprowadzone dla dwóch okresów:

– okres 1 roku – od lutego 2004 do lutego 2005 (18 funduszy) – okres ten był pierwszym rokiem funkcjonowania Polski jako kraju członkowskiego Unii Europejskiej, ponadto zakończono w tym okresie prace legislacyjne nad ustawą o funduszach inwestycyjnych, której najważniejszą częścią były regulacje przystosowujące polskie prawodawstwo do prawnych norm UE,

– okres 4 lat – od lutego 2001 do lutego 2005 (15 funduszy) – okres ten stanowi próbę oceny wyników osiąganych przez fundusze w długim okresie czasu, cechą charakterystyczną tego okresu była świadomość, iż znacznie zmienia się obowiązujące na rynku kapitałowym regulacje prawne. Dodatkowo w sposób istotny na kształt ówczesnego rynku funduszy wpływały takie wydarzenia, jak częste zmiany stóp procentowych czy wprowadzenie podatku dochodowego tzw. podatku Belki.

W oparciu o szeregi dziennych jednostek rozrachunkowych, obliczone zostały logarytmiczne tygodniowe stopy zwrotu, które kolejno wykorzystano do skonstruowania nakładów i wyników wykorzystanych w ocenie ich efektywności.

Zmienne wykorzystane w modelu zostały podzielone na dwie grupy:

Nakłady – jako nakłady w przeprowadzonym badaniu potraktowano: odchylenie standardowe stóp zwrotu funduszy, parametr β jako miarę ryzyka systematycznego, prowizje pobierane przez fundusze przy zakupie jednostek uczestnictwa.

³ Model został wykorzystany do oceny efektywności rynku funduszy we Włoszech [1].

Wyniki – jako pierwszą zmienną w grupie wyniki wykorzystano stopę zwrotu osiągniętą przez badane fundusze odpowiednio w okresie 1 roku i 4 lat, natomiast drugą zmienną stanowił współczynnik asymetrii stóp zwrotu – w okresie 1 roku i 4 lat.

W procesie wyznaczania funduszy efektywnych w badanej grupie funduszy akcyjnych optymalizowanych było sześć wariantów modelu zorientowanego na nakłady, różniące się między sobą zestawem zmiennych stanowiących nakłady i wyniki. Dodatkowo każdy z sześciu wariantów optymalizowany był zarówno dla okresu rocznego, jak i czteroletniego. Zestawy zmiennych w kolejnych wariantach optymalizowanego modelu DEA były następujące:

– DEA_{11c} – jeden nakład (odchylenie standardowe stóp zwrotu) i jeden wynik (stopa zwrotu),

– DEA_{11b} – jeden nakład (parametr β stóp zwrotu) i jeden wynik (stopa zwrotu),

– DEA_{21c} – dwa nakłady (odchylenie standardowe stóp zwrotu i prowizje) i jeden wynik (stopa zwrotu),

– DEA_{21b} – dwa nakłady (parametr β stóp zwrotu i prowizje) i jeden wynik (stopa zwrotu),

– DEA_{22c} – dwa nakłady (odchylenie standardowe stóp zwrotu i prowizje) i dwa wyniki (stopa zwrotu i współczynnik asymetrii stóp),

– DEA_{22b} – dwa nakłady (parametr β stóp zwrotu i prowizje) i dwa wyniki (stopa zwrotu i współczynnik asymetrii stóp zwrotu).

W tabelach 1-3 zamieszczono wartości współczynników efektywności dla wszystkich rozważanych funduszy akcyjnych.

Tabela 1

Wartości współczynników efektywności oraz ranking dla wariantu DEA₁₁

Nazwa funduszu	Ryzyko całkowite				Ryzyko systematyczne			
	1 rok		4 lata		1 rok		4 lata	
Arka BZ WBK Akcji	0,798	7	0,49	11	0,587	7	0,393	11
CA IB FIO Akcji	0,624	13	0,547	9	0,44	14	0,472	8
DWS Polska FIO Akcji	0,672	10	0,388	15	0,504	11	0,303	15
DWS Polska FIO Akcji Plus	1	1	0,557	8	0,698	2	0,453	9
ING Akcji	0,548	15	0,481	13	0,387	16	0,367	13
CitiAkcji A	0,705	8	0,51	10	0,512	10	0,415	10
UniKorona Akcji	0,932	2	0,658	4	0,663	3	0,507	7
Pioneer Akcji Polskich	0,531	16	0,436	14	0,408	15	0,351	14
PKO/CS Akcji	0,831	5	0,622	6	0,634	4	0,519	6
Krakowiak	0,676	9	0,671	3	0,495	12	0,538	4

cd. tabeli 1

Nazwa funduszu	Ryzyko całkowite				Ryzyko systematyczne			
	1 rok		4 lata		1 rok		4 lata	
SEB 3	0,642	12	0,483	12	0,525	9	0,385	12
Skarbiec Akcja	0,647	11	0,643	5	0,458	13	0,528	5
Millenium FIO Akcji	0,407	17	0,563	7	0,316	18	0,611	3
CU Polskich Akcji	0,832	4	1	1	0,617	6	0,976	2
GTFI Akcji	0,591	14	–	–	0,373	17	–	–
Allianz Akcji	0,824	6	–	–	0,633	5	–	–
AIG Akcji	0,898	3	–	–	1	1	–	–
DWS Top 25	0,314	18	0,815	2	0,526	8	1	1

Źródło: opracowanie własne.

Dla pierwszego wariantu metody DEA_{11c} jedynie fundusz DWS Akcji Plus był efektywny w okresie jednego roku, natomiast w okresie czterech lat był to fundusz CU Polskich Akcji. W przypadku uwzględnienia jako nakładu ryzyka systematycznego mierzonego współczynnikiem β efektywny był fundusz AIG Akcji w okresie jednego roku, natomiast dla czterech lat był to fundusz DWS Top 25. Porównując miejsca w rankingu poszczególnych funduszy, można zauważyć zgodność w przypadku różnych miar ryzyka dla tego samego okresu analizy (jeden rok lub cztery lata). Natomiast porównując klasyfikację dla tej samej miary ryzyka, ale różnych okresów, to można zauważyć pewne różnice, które mogą wynikać z różnej liczebności grup. Jest to szczególnie widoczne w przypadku funduszu DWS Akcji Plus, który w okresie jednego roku zajmował pierwsze miejsce, natomiast dla okresu czterech lat było to miejsce 8.

Tabela 2

Wartości współczynników efektywności oraz ranking dla wariantu DEA₂₁

Nazwa funduszu	Ryzyko całkowite				Ryzyko systematyczne			
	1 rok		4 lata		1 rok		4 lata	
Arka BZ WBK Akcji	1	1	0,815	2	1	1	0,815	3
CA IB FIO Akcji	0,812	8	0,717	6	0,812	4	0,717	6
DWS Polska FIO Akcji	0,673	13	0,388	15	0,503	16	0,309	15
DWS Polska FIO Akcji Plus	1	1	0,557	11	0,726	5	0,463	13
ING Akcji	0,595	16	0,584	10	0,539	13	0,584	8
CitiAkcji A	0,749	9	0,512	12	0,637	8	0,512	11

cd. tabeli 2

Nazwa funduszu	Ryzyko całkowite				Ryzyko systematyczne			
	1 rok		4 lata		1 rok		4 lata	
UniKorona Akcji	1	1	0,767	5	0,874	3	0,767	5
Pioneer Akcji Polskich	0,538	17	0,436	14	0,408	18	0,36	14
PKO/CS Akcji	0,831	7	0,622	9	0,634	9	0,528	10
Krakowiak	0,676	12	0,671	8	0,495	17	0,549	9
SEB 3	0,682	11	0,504	13	0,585	12	0,504	12
Skarbiec Akcja	0,735	10	0,768	4	0,718	6	0,768	4
Millenium FIO Akcji	0,614	15	0,672	7	0,614	11	0,672	7
CU Polskich Akcji	0,839	6	1	1	0,623	10	1	1
GTFI Akcji	0,625	14	–	–	0,518	15	–	–
Allianz Akcji	0,85	5	–	–	0,669	7	–	–
AIG Akcji	1	1	–	–	1	1	–	–
DWS Top 25	0,314	18	0,815	2	0,526	14	1	1

Źródło: opracowanie własne.

W drugim wariancie metody DEA, aż cztery fundusze były efektywne w okresie jednego roku (dla nakładu ryzyka całkowitego mierzonego odchyleniem standardowym DEA_{21c}), natomiast w okresie czterech lat był to tylko jeden fundusz. Dla wariantu metody DEA_{21b} zarówno w okresie jednego roku, jak i dla czterech lat, efektywne były dwa fundusze. CU Polskich Akcji. Porównując miejsca w rankingu poszczególnych funduszy, można zauważyć zgodność w klasyfikacji, analogicznie jak w przypadku wariantu DEA_{11} . Zupełny brak zgodności występuje natomiast w przypadku funduszu DWS Top 25.

Tabela 3

Wartości współczynników efektywności oraz ranking dla wariantu DEA_{22}

Nazwa funduszu	Ryzyko całkowite				Ryzyko systematyczne			
	1 rok		4 lata		1 rok		4 lata	
Arka BZ WBK Akcji	1	1	1	1	1	1	1	1
CA IB FIO Akcji	0,988	7	1	1	0,977	5	1	1
DWS Polska FIO Akcji	0,821	14	0,774	11	0,775	13	0,626	14
DWS Polska FIO Akcji Plus	1	1	0,766	13	0,726	15	0,652	12
ING Akcji	0,826	12	0,791	9	0,786	12	0,75	8

cd. tabeli 3

Nazwa funduszu	Ryzyko całkowite				Ryzyko systematyczne			
	1 rok		4 lata		1 rok		4 lata	
CitiAkcji A	0,826	12	0,786	10	0,792	11	0,713	10
UniKorona Akcji	1	1	0,864	8	0,974	6	0,834	7
Pioneer Akcji Polskich	0,63	16	0,637	15	0,601	16	0,54	15
PKO/CS Akcji	1	1	0,741	14	1	1	0,634	13
Krakowiak	0,698	15	0,874	7	0,506	18	0,73	9
SEB 3	0,84	11	0,768	12	0,849	8	0,699	11
Skarbiec Akcja	0,883	9	0,89	6	0,864	7	0,87	6
Millenium FIO Akcji	1	1	0,954	5	1	1	1	1
CU Polskich Akcji	0,884	8	1	1	0,728	14	1	1
GTFI Akcji	0,625	17	–	–	0,518	17	–	–
Allianz Akcji	0,883	9	–	–	0,799	10	–	–
AIG Akcji	1	1	–	–	1	1	–	–
DWS Top 25	0,38	18	1	1	0,833	9	1	1

Źródło: opracowanie własne.

Dla trzeciego wariantu metody DEA otrzymano największą liczbę funduszy efektywnych, były to od czterech do sześciu funduszy w zależności od stosowanego wariantu metody DEA. Tak duża liczba podmiotów efektywnych w przypadku trzeciego modelu DEA wynika z faktu, że wraz ze wzrostem liczby nakładów i wyników, które są uwzględnione w modelu, rośnie liczba podmiotów efektywnych [3]. Analizując otrzymane wyniki wskaźników efektywności dla okresu jednego roku na pierwszych miejscach znajdowały się następujące fundusze: Arka Akcji, DWS Akcji Plus, UniKorona Akcji, PKO/CS Akcji, AIG Akcji. Natomiast w przypadku okresu czterech lat, w grupie najlepszych funduszy znalazły się ponownie Arka Akcji i UniKorona Akcji, a ponadto CU Polskich Akcji oraz DWS Top 25.

W kolejnym etapie przeprowadzonych badań porównano wyniki otrzymane za pomocą metody DEA z klasycznymi wskaźnikami, takimi jak: wskaźnik Jensena⁴, Treynora⁵ i Sharpe'a⁶. W tym celu dla każdego z funduszy obliczono wartość tych

⁴ Wskaźnik Jensena obliczono zgodnie ze wzorem: $J_j = (r_j - r_f) - \beta_j(r_m - r_f)$.

⁵ Wskaźnik Treynora obliczono zgodnie ze wzorem: $T_j = \frac{r_j - r_f}{\beta_j}$.

⁶ Wskaźnik Sharpe'a obliczono zgodnie ze wzorem: $S_j = \frac{r_j - r_f}{\sigma_j}$.

współczynników dla obu rozważanych okresów. Do obliczania tych klasycznych wskaźników potrzebne są takie elementy dodatkowe, jak: stopa zwrotu waloru wolnego od ryzyka (r_f), stopa zwrotu (r_m) oraz ryzyko portfela rynkowego (σ_m). W przeprowadzonym badaniu jako walor wolny od ryzyka wykorzystano średni ważony zysk z bonów skarbowych, natomiast indeks giełdowy WIG20 potraktowano, jako portfel rynkowy. W celu sprawdzenia, czy metoda DEA właściwie klasyfikuje fundusze obliczono współczynniki korelacji między współczynnikami efektywności w kolejnych wariantach metody DEA a wartościami wskaźników klasycznych. Wyniki otrzymane dla obu okresów (jednego roku i czterech lat) zamieszczono w poniższych tabelach 4 i 5.

Tabela 4

Współczynniki korelacji między wynikami DEA i wskaźnikami klasycznymi w okresie jednego roku

Wartości wskaźników klasycznych	Wartości współczynników efektywności otrzymane dla metody DEA					
	DEA _{11c}	DEA _{11b}	DEA _{21c}	DEA _{21b}	DEA _{22c}	DEA _{22b}
Wskaźnik Jensena	0,88	0,87	0,88	0,74	0,65	0,34
Wskaźnik Treynora	0,89	0,83	0,88	0,67	0,70	0,29
Wskaźnik Sharpe'a	0,98	0,71	0,95	0,63	0,74	0,21

Źródło: opracowanie własne.

Dla okresu jednego roku wartości współczynników korelacji są bardzo wysokie w przypadku wariantów metody DEA_{11c}, DEA_{11b}, DEA_{21c}, czyli gdy występuje wariant jednego nakładu i jednego wyniku, oraz dwóch nakładów i jednego wyniku, a nakładem są odchylenie standardowe i koszty. Otrzymane wyniki wskazują na dużą zgodność w otrzymanych wynikach dla tych dwóch różnych metod oceny efektywności zarządzania portfelem inwestycji. Niska wartość współczynników korelacji występuje jedynie w przypadku ostatniego wariantu metody DEA, czyli tego, w którym nakłady tworzą parametr β i koszty, natomiast wyniki to stopa zwrotu oraz współczynnik asymetrii stóp zwrotu.

Tabela 5

Współczynniki korelacji między wynikami DEA i wskaźnikami klasycznymi w okresie czterech lat

Wartości wskaźników klasycznych	Wartości współczynników efektywności otrzymane dla metody DEA					
	DEA _{11c}	DEA _{11b}	DEA _{21c}	DEA _{21b}	DEA _{22c}	DEA _{22b}
Wskaźnik Jensena	0,99	0,93	0,86	0,83	0,60	0,59
Wskaźnik Treynora	0,94	1,00	0,78	0,81	0,62	0,61
Wskaźnik Sharpe'a	1,00	0,92	0,84	0,79	0,56	0,53

Źródło: opracowanie własne.

Dla okresu czterech lat wartości współczynników korelacji wzrosły w większości przypadków dla wariantów pierwszego (DEA_{11}) i drugiego (DEA_{22}), w przypadkach zależności między wskaźnikiem Sharpe'a oraz DEA_{11c} oraz wskaźnikiem Treynora a DEA_{11b} wartości współczynnika korelacji wskazuje na istnienie dodatniej zależności liniowej. Analizując wartości współczynników korelacji dla trzeciego wariantu metody DEA, można zauważyć spadek wartości współczynników w stosunku do okresu jednego roku dla wariantu DEA_{22c} oraz znaczny wzrost w przypadku DEA_{22b} .

5. PODSUMOWANIE

Metoda DEA jest jedną z technik badań operacyjnych, która szeroko stosowana w ocenie efektywności działalności gospodarowania. W pracy podjęto próbę zastosowania metody DEA w ocenie efektywności wyników osiąganych przez fundusz inwestycyjny. Otrzymane wyniki badań pokazały, że metoda DEA jest narzędziem, które może być wykorzystane w procesie podejmowania decyzji inwestycyjnych na rynku kapitałowym.

W świetle przeprowadzonych badań można stwierdzić, że grupa akcyjnych funduszy inwestycyjnych jest zróżnicowana pod względem poziomu efektywności. Zastosowana metoda pozwoliła na uwzględnienie wielu czynników determinujących opłacalność inwestycji w fundusz, takich jak: stopa zwrotu funduszu, ryzyko funduszu, koszt uczestnictwa oraz horyzont czasowy inwestycji. Przeprowadzona analiza porównawcza z klasycznymi wskaźnikami oceny efektywności zarządzania portfelem inwestycyjnym, takimi jak wskaźniki: Treynora, Sharpe'a, Jensena, wskazuje na zgodność uzyskanych klasyfikacji funduszy w oparciu o te wskaźniki z klasyfikacją uzyskaną za pomocą metody DEA. Zgodność tych klasyfikacji wskazuje, że metoda DEA może być stosowana jako alternatywa albo uzupełnienie klasycznych wskaźników w ocenie efektywności zarządzania portfelem inwestycyjnym funduszu.

Uniwersytet Gdański

LITERATURA

- [1] Basso A., Funari S., [2001], *A data envelopment analysis approach to measure the mutual fund performance*, „European Journal of Operational Research”, 135, s. 477-492.
- [2] Charnes A., Cooper W.W., Rhodes E., [1978], *Measuring the efficiency of decision making units*, „European Journal of Operational Research”, 3, 429-444.
- [3] Galagedera D., Silvapulle P., [2002], *A Australian mutual fund performance appraisal using data envelopment analysis*, Managerial Finance, 9, 60-73.
- [4] Jensen M., [1968], *The performance of mutual funds in the period 1945-1964*, „Journal of Finance”, 50, 549-572.
- [5] Murthi B., Choi Y., Desai P., [1997], *Efficiency of Mutual Funds and Portfolio Performance Measurement: A Non-Parametric Approach*, „European Journal of Operational Research”, 98, 408-418.
- [6] Osiewalska A., Osiewalski J., [1999], *Próba oceny efektywności kosztowej polskich bibliotek akademickich*, EBIB, <http://www.oss.wroc.pl/biuletyn/ebib03/efektywn.html>

- [7] Pawłowska M., [2005], *Konkurencja i efektywność na polskim rynku bankowym na tle zmian strukturalnych i technologicznych*, Materiały i Studia NBP, Zeszyt nr 192.
- [8] Prędko A., [2003], *Analiza efektywności za pomocą metody DEA: podstawy formalne i ilustracja ekonomiczna*, „Przegląd Statystyczny”, 1, 87-100.
- [9] Ruiyue L., Zhiping C., [2008], *New Dea Performance Evaluation Indices and Their Application in the American Fund Market*, Asia – Pacific Journal of Operational Research, 4, 421-450.
- [10] Seiford L., Thrall R., [1990], *Recent Developments in DEA*, „Journal of Econometrics”, 46, 7-38.
- [11] Sharpe W.F., [1994], *The Sharpe Ratio*, „Journal of Portfolio Management”, 49-59.
- [12] Sharpe W.F., [1966], *Mutual Fund Performance*, „Journal of Business”, 34, 119-138.
- [13] Treynor J.L., [1965], *How to rate management investment funds*, Harvard Business Review, 43, 63-75.
- [14] Yun Y.B., Nakayama H., Tanino T., [2004], *A generalized model for data envelopment analysis*, „European Journal of Operational Research”, 157, 87-105.
- [15] Zamojska A., [2006], *Zastosowanie metody DEA w ocenie efektywności zarządzania portfelem funduszu*, Taksonomia 13, 112-120.

Praca wpłynęła do redakcji w październiku 2009 r.

ZASTOSOWANIE METODY DEA W KLASYFIKACJI FUNDUSZY INWESTYCYJNYCH

Streszczenie

Artykuł prezentuje zastosowanie metody DEA do mierzenia efektywności zarządzania portfelem inwestycyjnym funduszy akcyjnych funkcjonujących na polskim rynku kapitałowym. Do oceny efektywności badanych funduszy wykorzystano sześć wariantów metody DEA, różniących się między sobą strukturą wyników i nakładów i w oparciu o otrzymane wyniki sporządzono klasyfikację funduszy. Kolejno sporządzono klasyfikację funduszy w oparciu o klasyczne wskaźniki: Treynora, Sharpe'a, Jensena i porównano obie klasyfikacje. Przedstawione wyniki badań wskazują, że metoda DEA może być stosowana jako alternatywa albo uzupełnienie klasycznych wskaźników w ocenie efektywności zarządzania portfelem inwestycyjnym funduszu i stanowi użyteczne narzędzie, które może być wykorzystywane w procesie podejmowania decyzji inwestycyjnych.

Słowa kluczowe: DEA, Efektywność, Fundusze Inwestycyjne

EFFICIENCY OF MUTUAL FUND AND A DATA ENVELOPMENT ANALYSIS

Summary

The paper presents an application of the DEA method to measure the efficiency of equity mutual funds on the Polish capital market. The DEA method is an alternative measure of performance of the funds. The advantage of the DEA is that this particular method does not require any assumption concerning distributions of returns. Results showed that the DEA is a simple and very useful at the same time tool for creating the ranks of funds.

Key words: DEA, Efficiency, Mutual Fund

ANDRZEJ STRYJEK

ZASTOSOWANIE MIAR ZALEŻNOŚCI ZMIENNYCH LOSOWYCH ORAZ KOPULI CLAYTONA I GUMBEL-HOUGAARDA DO SZACOWANIA WARTOŚCI ZAGROŻONEJ¹

Od początku lat dziewięćdziesiątych ubiegłego wieku istotną rolę w ocenie ryzyka finansowego odgrywa wartość zagrożona (Value at Risk – VaR). Regulacje prawne w Polsce nakładają na instytucje finansowe obowiązek stosowania tej miary w celu bieżącej oceny poziomu ryzyka portfela, lecz jednocześnie pozostawiają dużą dowolność w wyborze algorytmu estymacji. Stale prowadzone są badania nad różnymi metodami szacowania wartości zagrożonej. Bieżący artykuł ma na celu zaprezentowanie możliwości, jakie daje zastosowanie kopuli do wyznaczania wartości VaR oraz porównanie tej metody z innymi, bardziej znanymi.

Kopula jest pojęciem mocno osadzonym w teorii statystyki. W ostatnich latach pojęcie to przeżywa swój renesans z powodu licznych prób zastosowania do modelowania zależności zmiennych opisujących szeregi finansowe. Zagadnienie szacowania wartości zagrożonej za pomocą kopuli pojawia się w pracach [2] i [5].

W opracowaniu oprócz omówienia pojęcia wartości zagrożonej oraz funkcji powiązań (kopuli) autor prezentuje własną, nową metodę szacowania VaR za pomocą kopuli Clayтона oraz kopuli Gumbel-Hougaard, a także wyniki badania empirycznego ilustrującego efektywność różnych metod wyznaczania VaR za pomocą kopuli.

1. WARTOŚĆ ZAGROŻONA

Wartość zagrożona VaR jest definiowana jako maksymalna możliwa strata wartości rynkowej portfela, której przekroczenie w rozpatrywanym horyzoncie czasowym inwestycji może nastąpić z zadany, niewielkim prawdopodobieństwem $\alpha \in (0, 1)$ zwanym poziomem tolerancji. Decyzja o długości okresu na jaki wyznacza się VaR oraz poziomie tolerancji α jest w gestii podmiotu dokonującego inwestycji. Uściślając, wartość zagrożona jest określona jako liczba $VaR_\alpha(X)$, która spełnia warunek

$$VaR_\alpha(X) = -\sup\{x \in \mathbb{R}: P(X \leq x) \leq \alpha\},$$

¹ Opracowanie jest kontynuacją badań przeprowadzonych w badaniu statutowym nr 03/S/0015/07 w Instytucie Ekonometrii Szkoły Głównej Handlowej w Warszawie (zob. [6]).

gdzie parametr $\alpha \in (0, 1)$ oznacza poziom tolerancji, natomiast X jest zmienną losową opisującą potencjalne zyski i straty, które może wygenerować rozpatrywany portfel. Zmienna ta nazywana jest zmienną zysków i strat lub zmienną P&L².

Ponieważ $VaR_\alpha(X)$ jest górnym kwantylem rozkładu zmiennej P&L, więc w przypadku modelowania zmiennej X za pomocą ciągłego rozkładu prawdopodobieństwa, wartość $VaR_\alpha(X)$ jest α -kwantylem tego rozkładu. Można zatem zapisać, że $VaR_\alpha(X) = -F^{-1}(\alpha)$, przy czym F jest dystrybuantą zmiennej X .

Przy założeniu, że X jest zmienną o normalnym rozkładzie prawdopodobieństwa ze średnią $\mu \in R$ oraz odchyleniem standardowym $\sigma > 0$ otrzymujemy

$$VaR_\alpha(X) = -\phi_{\mu,\sigma}^{-1}(\alpha) = -\mu - \phi^{-1}(\alpha) \cdot \sigma,$$

gdzie $\phi_{\mu,\sigma}$ jest dystrybuantą wspomnianego rozkładu normalnego oraz $\phi = \phi_{0,1}$.

W praktyce działania instytucji finansowej pojawia się zagadnienie szacowania wartości zagrożonej dla portfela instrumentów finansowych. Przyjęcie założenia, że zmienne P&L poszczególnych instrumentów wchodzących w skład portfela mają rozkład normalny upraszcza proces wyznaczenia wartości VaR dla portfela. W takim przypadku portfel złożony z k inwestycji, które opisane są zmiennymi X_i dla $i = 1, 2, \dots, k$ oraz proporcjach składników portfela $\varphi_i \geq 0$ dla $i = 1, 2, \dots, k$, gdzie $\varphi_1 + \varphi_2 + \dots + \varphi_k = 1$ jest określony przez zmienną zysków i strat $X_{P\&L}$. Rozkład tej zmiennej jest także rozkładem normalnym. Wartość zagrożona dla portfela wynosi

$$VaR_\alpha(X_{P\&L}) = -\sum_{i=1}^k \varphi_i \cdot \mu_i - \phi^{-1}(\alpha) \cdot \sqrt{\sum_{i=1}^k (\varphi_i^2 \cdot \sigma_i^2) + 2 \sum_{i < j} (\varphi_i \cdot \varphi_j \cdot \sigma_{ij})},$$

gdzie

μ_i – wartość oczekiwana zmiennej X_i ,

σ_i^2 – wariancja zmiennej X_i ,

σ_{ij} – kowariancja zmiennych X_i, X_j dla $i \neq j$.

Szacowanie wartości zagrożonej za pomocą powyższego wzoru, w literaturze przedmiotu, nosi nazwę metody kowariancji. Metoda ta jest prosta w zastosowaniu, lecz założenie o normalnym rozkładzie prawdopodobieństwa zmiennej P&L często stanowi nadmierne uproszczenie. Drugą metodą szacowania VaR jest metoda historyczna. Polega na uszeregowaniu historycznych realizacji zmiennej X , a następnie odczytaniu empirycznego kwantyla rzędu α . Metoda ta jest nieparametryczna i nie wymaga przyjmowania żadnych założeń o rozkładzie prawdopodobieństwa zmiennej zysków i strat. Niestety, dużą jej wadą jest fakt, że do przewidywania kształtowania się możliwych zmian wartości portfela korzysta się tylko z danych historycznych i często przy jej zastosowaniu nie jest spełniony warunek dywersyfikacji ryzyka (por. [8]). Wreszcie trzecią metodą szacowania VaR jest symulacja Monte Carlo. Polega na wielokrotnym wygenerowaniu obserwacji z rozkładu zmiennej X i wyznaczeniu za każdym razem kwantyla rzędu α . Ostatecznie za ocenę wartości VaR przyjmuje się średnią z otrzymana-

² Niektórzy autorzy (patrz [3] oraz [5]) zamiast zmiennej P&L posługują się równoważnie stopą zwrotu inwestycji R_t .

nych kwantyli we wszystkich iteracjach. Alternatywnie, zamiast generować dużą liczbę próbek, można wygenerować jedną, dużą próbę z zadanego rozkładu i odczytać VaR jako odpowiedni kwantyl. Tu w obliczaniu wartości zagrożonej znajdują zastosowanie także inne rozkłady niż rozkład normalny np. *t*-studenta, GED, rozkłady α -stabilne, mieszane rozkłady normalne.

Wyboru metody obliczania wartości zagrożonej dokonuje się na podstawie porównania w wybranym okresie wszystkich wyznaczonych wartości VaR z rzeczywistymi zmianami wartości inwestycji. Są to tzw. testy wsteczne. Najprostsza możliwość to wyznaczenie udziału przekroczeń rzeczywistej zmiany wartości portfela ponad wyznaczone wartości VaR w całej próbie testowej. Udział ten nie powinien przekraczać przyjętego poziomu α . Drugim sposobem jest zastosowanie testu Kupca (1995). Hipoteza zerowa tego testu o poprawności metody szacowania VaR jest weryfikowana za pomocą statystyki:

$$LR_{uc} = -2 \ln((1 - \alpha)^{T-N} \alpha^N) + 2 \ln\left(\left(1 - \frac{N}{T}\right)^{T-N} \left(\frac{N}{T}\right)^N\right),$$

gdzie

- α – jest przyjętym poziomem tolerancji w testowanym modelu,
- T – oznacza długość próby testowej,
- N – liczbę zaobserwowanych przekroczeń w tej próbie.

Statystyka LR_{uc} ma asymptotyczny rozkład χ^2 z jednym stopniem swobody. Odrzucenie hipotezy zerowej przy poziomie istotności testu wynoszącym 0,05 następuje, gdy wartość statystyki przekracza 3,84.

2. KOPULA

W ostatnich latach dzięki wykorzystaniu mocy obliczeń komputerowych pojęcie kopuli próbuje się coraz częściej stosować do modelowania zależności zmiennych opisujących rynki finansowe.

Kopula jest to funkcja $C: I^n \rightarrow \mathbb{R}$, gdzie $I = [0, 1]$, $n \geq 2$, która spełnia następujące warunki:

- i. $C(u_1, u_2, \dots, u_{k-1}, 0, u_{k+1}, \dots, u_n) = 0$ dla dowolnych $u_i \in [0, 1]$ oraz $i = 1, 2, \dots, k-1, k+1, \dots, n$,
- ii. $C(1, \dots, 1, u_i, 1, \dots, 1) = u_i$ dla dowolnego $u_i \in [0, 1]$ oraz $i = 1, 2, \dots, n$,
- iii. $V_C([\mathbf{u}, \mathbf{v}]) = \sum_{\mathbf{w}} \text{sgn}(\mathbf{w}) C(\mathbf{w}) \geq 0$,

gdzie:

$[\mathbf{u}, \mathbf{v}] = [u_1, v_1] \times \dots \times [u_n, v_n]$, przy czym $u_i \leq v_i$ dla $i \in \{1, 2, \dots, n\}$,

$\mathbf{u} = (u_1, u_2, \dots, u_n) \in I^n$, $\mathbf{v} = (v_1, v_2, \dots, v_n) \in I^n$,

sumowanie przebiega po wszystkich wierzchołkach $\mathbf{w} = (w_1, w_2, \dots, w_n)$ n -wymiarowej kostki $[\mathbf{u}, \mathbf{v}]$,

$$\text{sgn}(\mathbf{w}) = \begin{cases} +1, & \text{jeżeli } w_k = u_k \text{ dla parzystej liczby współrzędnych,} \\ -1, & \text{jeżeli } w_k = u_k \text{ dla nieparzystej liczby współrzędnych.} \end{cases}$$

W klasie rozkładów eliptycznych standardową miarą zależności zmiennych losowych jest współczynnik korelacji liniowej Pearsona. Dla zmiennych spoza tej klasy pomiar stopnia zależności za pomocą tego współczynnika może prowadzić do błędnych wniosków.

Pojęcie kopuli jest stosowane do ustalania stopnia zależności między zmiennymi losowymi o dowolnych rozkładach. Co więcej, dla dowolnych ściśle rosnących funkcji f i g , jeżeli zmiennym X i Y odpowiada kopula C , to zmiennym $f(X)$ oraz $g(Y)$ także odpowiada kopula C . W przypadku zastosowania współczynnika Pearsona taka własność zachodzi tylko dla rosnących funkcji liniowych.

Bezpośrednią możliwość wykorzystania kopuli do modelowania zmiennych na rynkach finansowych stworzył A. Sklar formułując następującą własność³:

Niech funkcje $F_1(x)$ oraz $F_2(y)$ będą dystrybuantami jednowymiarowych zmiennych losowych X i Y . Dla dowolnych $(x, y) \in \mathbb{R}^2$ zachodzi⁴:

- jeżeli funkcja C jest kopulą, to $C(F_1(x), F_2(y))$ jest dystrybuantą łącznego rozkładu wektora (X, Y) o rozkładach brzegowych danych dystrybuantami $F_1(x)$ oraz $F_2(y)$,
- jeżeli $F(x, y)$ jest dystrybuantą łącznego rozkładu wektora (X, Y) o rozkładach brzegowych zadanych dystrybuantami $F_1(x)$ oraz $F_2(y)$, to istnieje dokładnie jedna kopula C taka, że $F(x, y) = C(F_1(x), F_2(y))$.

Z twierdzenia Sklara wynikają liczne zastosowania kopuli. Po pierwsze, znając rozkłady brzegowe i łączny rozkład prawdopodobieństwa wektora zmiennych losowych można dopasować odpowiednią kopulę. Wyznaczona kopula może być zastosowana do wyznaczania miar zależności między zmiennymi np. współczynnika korelacji liniowej Pearsona, τ Kendalla, ρ_S Spearmana.

Po drugie, kopula umożliwia dokonanie symulacji łącznego rozkładu prawdopodobieństwa wektora zmiennych przy zadanych rozkładach brzegowych. Oczywiście, w praktycznych zagadnieniach etap symulacji musi być poprzedzony oszacowaniem nieznanymi parametrów kopuli na podstawie danych empirycznych.

Najczęściej stosowaną metodą estymacji parametrów kopuli jest estymacja metodą największej wiarygodności. Do zastosowania tej metody niezbędna jest znajomość parametrów rozkładów brzegowych, co w praktyce oznacza konieczność dokonania ich estymacji innymi metodami. Dlatego też w literaturze funkcjonuje inny sposób estymacji parametrów kopuli tzw. wnioskowanie dla rozkładów brzegowych. Ta metoda oparta jest na spostrzeżeniu, że w zlogarytmowanej funkcji wiarygodności występują dwa składniki, przy czym jeden zawiera tylko parametry rozkładów brzegowych, a drugi zawiera dodatkowo także parametry kopuli. Estymacja odbywa się wówczas w dwóch krokach. W pierwszym szacuje się parametry rozkładów brzegowych poprzez ich dobór tak, aby maksymalizować pierwszy składnik. W drugim kroku dokonuje się doboru parametrów kopuli, aby zmaksymalizować drugi składnik przy wyznaczonych estyma-

³ Sklar A. (1959), *Fonctions de repartition a n dimensions et leurs margers*, Pub. Inst. Statist. Univ. Paris, 8, 229-231.

⁴ $\mathbb{R} = \mathbb{R} \cup \{-\infty, +\infty\}$.

torach z pierwszego kroku. Obie procedury są w praktyce czasochłonne. W przypadku kopuli zależnej od jednego parametru można podać prostszy sposób estymacji tego parametru. Polega na wyznaczeniu oszacowania współczynnika τ Kendalla (lub ρ_S Spearmana) na podstawie danych. W dalszym ciągu procedury wyznaczenie oszacowania parametru kopuli sprowadza się do rozwiązywania nieskomplikowanego równania (por. [2], [4]).

Wiele przykładów kopuli jest zawartych w literaturze przedmiotu (zob. [3]). W badaniu empirycznym, którego wyniki są zaprezentowane w dalszej części artykułu, wykorzystano trzy funkcje powiązań:

(1) Kopula Gaussa

$$C(u, v; \theta) = \int_{-\infty}^{\phi^{-1}(u)} \int_{-\infty}^{\phi^{-1}(v)} \frac{1}{2\pi\sqrt{1-\theta^2}} \exp\left(\frac{-(s^2 - 2\theta \cdot st + t^2)}{2(1-\theta^2)}\right) ds dt,$$

gdzie ϕ jest dystrybuantą standardowego rozkładu normalnego oraz $\theta \in [-1, 1]$.

(2) Kopula Claytona

$$C(u, v; \theta) = (u^{-\theta} + v^{-\theta} - 1)^{\frac{1}{\theta}}, \theta \in (0, +\infty)$$

(3) Kopula Gumbel-Hougaard

$$C(u, v; \theta) = \exp\left(-((-\ln u)^\theta + (-\ln v)^\theta)^{\frac{1}{\theta}}\right), \theta \in (0, +\infty).$$

Wybór tych funkcji jest arbitralną decyzją autora. Po części wynikał z faktu, że stanowią one jednoparametrowe rodziny funkcji, co w przypadku estymacji parametrów jest zagadnieniem istotnie prostszym niż dla kopuli o większej liczbie parametrów. W przypadku kopuli Claytona i Gumbel-Hougaard zaletą jest też nieskomplikowana postać analityczna. Pozostaje nadal otwartym pytanie o to, w jaki sposób spośród szerokiego spektrum funkcji powiązań dokonywać wyboru odpowiedniej kopuli do badanego zagadnienia.

3. OPIS ZASTOSOWANYCH PROCEDUR

Szczegółowy opis badania empirycznego zawarty jest w następnym punkcie artykułu. Bieżący punkt jest poświęcony zaprezentowaniu procedur, które zostały użyte przez autora we wspomnianej analizie empirycznej.

Wektory $\mathbf{X} = (x_1, x_2, \dots, x_n)$ oraz $\mathbf{Y} = (y_1, y_2, \dots, y_n)$ oznaczają dzienne zmiany wartości, odpowiednio, dwóch instrumentów finansowych obserwowanych w tym samym okresie. Na podstawie wartości zawartych w wektorach $\mathbf{X}$ i $\mathbf{Y}$ dokonano oszacowania wartości zagrożonej $Var_\alpha(\beta \cdot X + (1 - \beta) \cdot Y)$ gdzie $\beta = 0,1; 0,2; \dots; 0,9$, następującymi metodami: klasyczną metodą kowariancji, dwiema modyfikacjami tej metody, metodą

symulacji Monte Carlo oraz symulacji z zastosowaniem kopuli Clayтона i Gumbel-Hougaard.

Estymacja metodą kowariancji była przeprowadzona zgodnie ze wzorem (1). Łatwo zauważyć, że formuła (1) może zostać zapisana w postaci

$$VaR_{\alpha}(X_{P\&L}) = - \sum_{i=1}^k \varphi_i \cdot \mu_i - \phi^{-1}(\alpha) \cdot \sqrt{\sum_{i=1}^k (\varphi_i^2 \cdot \sigma_i^2) + 2 \sum_{i < j} (\varphi_i \cdot \varphi_j \cdot \sigma_i \cdot \sigma_j \cdot \rho_{ij})},$$

gdzie ρ_{ij} oznacza współczynnik korelacji między zmiennymi X_i oraz X_j .

Modyfikacja tej metody, dokonana przez autora, polegała na zastąpieniu współczynnika korelacji liniowej $\rho_{\mathbf{XY}}$ przez współczynnik τ Kendalla lub ρ_S Spearmana. Przypomnijmy, że

$$\hat{\tau}_{\mathbf{XY}} = \frac{2}{n(n-1)} \sum_{i=1}^n \sum_{i < j} \text{sgn}((x_i - x_j) \cdot (y_i - y_j)) \quad (2)$$

oraz

$$\hat{\rho}_s = 1 - 6 \cdot \frac{(r_i - s_i)^2}{n(n^2 - 1)},$$

gdzie r_i jest rangą obserwacji x_i , a s_i rangą obserwacji y_i .

Procedura jednej iteracji w symulacji Monte Carlo dla rozkładu normalnego została przeprowadzona w następujących etapach:

1. wyznaczenie macierzy kowariancji C dla wektorów $\beta \cdot \mathbf{X}$ i $(1 - \beta) \cdot \mathbf{Y}$,
2. dekompozycja Choleskiego macierzy C do postaci $\mathbf{A} \cdot \mathbf{A}^T$, gdzie $\mathbf{A}$ jest macierzą trójkątną dolną,
3. wygenerowanie dwóch szeregów obserwacji $\mathbf{Z}_X$ oraz $\mathbf{Z}_Y$ – obu ze standardowego rozkładu normalnego,
4. dokonanie transformacji $\mathbf{S} = \mathbf{A} \cdot \mathbf{Z}_X$ oraz $\mathbf{T} = \mathbf{A} \cdot \mathbf{Z}_Y$,
5. wektor $(\mathbf{S}, \mathbf{T})$ ma rozkład normalny o macierzy kowariancji C . Szereg symulowanych wartości dziennych zmian portfela to wektor $\mathbf{T} + \mathbf{S}$, zatem wartość zagrożona $VaR_{\alpha}(\beta \cdot \mathbf{X} + (1 - \beta) \cdot \mathbf{Y})$ jest odpowiednim kwantylem w jego rozkładzie empirycznym.

W procedurze wyznaczania wartości zagrożonej za pomocą kopuli Clayтона i Gumbel-Hougaard można wyróżnić dwa etapy.

Pierwszy etap dotyczy oszacowania parametru kopuli. Kopule Clayтона i Gumbel-Hougaard należą do grupy kopuli archimedesowych, których postać funkcyjna jest zależna od jednego parametru θ . Estymacja tego parametru polegała na wyznaczeniu w pierwszej kolejności oszacowania współczynnika τ Kendalla na podstawie wzoru (2). Ponieważ dla obu kopuli współczynnik τ Kendalla wyraża się za pomocą parametru kopuli tzn. dla kopuli Clayтона

$$\tau = \frac{\theta}{\theta + 2}, \theta \in (0, +\infty),$$

kopuli Gumbel-Hougaarda zaś

$$\tau = \frac{\theta - 1}{\theta}, \theta \in (0, +\infty),$$

więc estymator $\hat{\theta}$ obliczano jako rozwiązanie odpowiedniego równania. Taki sposób estymacji pojawia się w [3] i [4].

Drugi etap procedury szacowania wartości zagrożonej za pomocą kopuli Claytona i Gumbel-Hougaarda polegał na generowaniu dwóch szeregów wartości z rozkładu o oszacowanej kopule. Na tym etapie posłużono się algorytmem opisanym w pracy [4]:

1. wygenerowanie niezależnie z rozkładu jednostajnego na odcinku $[0, 1]$ pary liczb (q, s) ,

2. wyznaczenie wartości $t = K_C^{-1}(q)$, gdzie $K_C(q) = t - \frac{\varphi(t)}{\varphi'(t)}$, φ jest generatorem kopuli⁵, a K_C^{-1} jest funkcją odwrotną w uogólnionym sensie⁶,

3. przyjęcie $u = \varphi^{-1}(s \cdot \varphi(t))$ oraz $v = \varphi^{-1}((1-s) \cdot \varphi(t))$,

4. para (u, v) stanowi element szukanej próby.

Wielokrotne zastosowanie tego algorytmu generuje dwa szeregi wartości z rozkładu o zadanej kopule. Są to wartości dystrybuant rozkładów brzegowych. W celu otrzymania dwóch szeregów symulujących zmiany wartości składników portfela dokonano jeszcze transformacji tych szeregów za pomocą uogólnionej odwrotności dystrybuanty empirycznej odpowiedniego rozkładu brzegowego⁷ zmiennej P&L. W ten sposób nie było konieczności przyjmowania założenia o rozkładzie zmiennej zysków i strat poszczególnych składników portfela. Jest to nowe podejście, odbiegające od prezentowanych w literaturze metod symulacji zmian wartości portfela za pomocą kopuli.

4. OPIS I WYNIKI BADANIA EMPIRYCZNEGO

Opisane w tym punkcie badanie empiryczne ma na celu zaprezentowanie możliwości zastosowania kopuli do szacowania wartości zagrożonej. W analizie dokonano empirycznego porównania efektywności znanych z literatury i zaproponowanych przez autora różnych metod szacowania wartości VaR z podstawową metodą estymacji wartości zagrożonej, tj. metodą kowariancji.

Badanie empiryczne zostało przeprowadzone dla trzech szeregów dziennych zmian wartości akcji wybranych spółek, które były notowane na Giełdzie Papierów Wartościowych w Warszawie. Były to spółki wchodzące w skład indeksu mWIG40 (dawniej MIDWIG). Indeks mWIG40 obejmuje średnie spółki notowane na giełdzie, a zatem nie są one instrumentami o największej płynności (WIG20) ani największej zmienności (sWIG80). Indeks ten jest notowany od 31 grudnia 1997 r., co umożliwiło

⁵ Generatorem kopuli Claytona jest funkcja $\varphi(t) = \frac{t^{-\alpha} - 1}{\alpha}$, a kopuli Gumbel-Hougaarda $\varphi(t) = (-\ln(t))^\alpha$.

⁶ $K_C^{-1}(y) = \inf\{x: K_C(x) \geq y\}$.

⁷ $F^{-1}(y) = \inf\{x: F(x) \geq y\}$.

wybranie szeregów danych odpowiednich pod względem liczby obserwacji i trendu. Z drugiej strony skład indeksu zmieniany jest co kwartał (w marcu, czerwcu, wrześniu i grudniu). Zatem do badania empirycznego wybrane zostało dziesięć spółek w trzech okresach: od 24 maja 2000 do 11 października 2002, od 30 listopada 2004 do 18 kwietnia 2007 oraz od 2 grudnia 2005 do 24 kwietnia 2008, które praktycznie przez cały podany okres były obecne w indeksie mWIG40. W przypadku trzeciego z wymienionych okresów takich spółek było więcej, lecz dla zachowania spójności procedur numerycznych i porównywalności otrzymanych wyników, do badania wylosowano spośród nich dokładnie 10 spółek. Lista spółek uczestniczących w badaniu zawiera tabela 1. Dane zostały pobrane z serwisu ISI Emerging Markets (www.securities.com).

Tabela 1

Spółki uczestniczące w badaniu empirycznym

24.05.2000 – 11.10.2002		30.11.2004 – 8.04.2007		2.12.2005 – 24.04.2008	
nazwa	kod spółki	nazwa	kod spółki	nazwa	kod spółki
AMICA	PLAMICA00010	COMARCH	PLCOMAR00012	AMREST	NL0000474351
BUDIMEX	PLBUDMX00013	ECHO	PLECHPS00019	BUDIMEX	PLBUDMX00013
CERSANIT	PLCRSNT00011	GETIN	PLGSPR000014	CCC	PLCCC0000016
ECHO	PLECHPS00019	GRAJEW	PLZPW0000017	DUDA	PLDUDA000016
ELBUDOWA	PLELTBD00017	HANDLOWY	PLBH00000012	ECHO	PLECHPS00019
IMPEXMETAL	PLIMPXM00019	INGBSK	PLBSK0000017	GETIN	PLGSPR000014
JUTRZENKA	PLJTRZN00011	KREDYT	PLKRDTB00011	INGBSK	PLBSK0000017
KREDYT	PLKRDTB00011	LPP	PLLP00000011	KREDYTB	PLKRDTB00011
PGF	PLMEDCS00015	MILLENIU	PLBIG0000016	MILLENIU	PLBIG0000016
RAFAKO	PLRAFAK00018	RAFAKO	PLRAFAK00018	SWIECIE	PLCELZA00018

Źródło: opracowanie własne.

W badaniu analizowano dwuskładnikowe portfele w trzech okresach wyszczególnionych w tabeli 1, odpowiednio, okres o trendzie stałym, okres o trendzie rosnącym i okres o trendzie początkowo rosnącym ze zmianą trendu na malejący. Dla każdych dwóch spółek tworzone portfele z wagami β i $1 - \beta$, gdzie $\beta = 0,1; 0,2; \dots; 0,9$. Obliczenia przeprowadzono za pomocą procedur napisanych przez autora w programie R 2.7.1.

Wartość zagrożona była szacowana na podstawie 125, 250 oraz 500 obserwacji dziennych zmian wartości poszczególnych akcji w każdym badanym szeregu. Otrzymane w ten sposób oszacowanie VaR porównywano z rzeczywistą zmianą wartości portfela w horyzoncie jednego dnia. Następnie obliczenia powtórzono dla szeregu przesuniętego o jeden dzień do przodu. Tę procedurę powtarzano sto razy, zatem okres testowy służący do oceny estymatorów wartości VaR obejmował sto dni.

W pierwszym etapie badania wykorzystano trzy wymienione wyżej szeregi danych w trzech wariantach długości szeregu (125, 250, 500 obserwacji). W tej części dokonano porównania klasycznej metody kowariancji z dwiema modyfikacjami tej metody opisanymi w poprzednim punkcie opracowania. Przypomnijmy, że modyfikacja metody kowariancji polegała na zastąpieniu współczynnika korelacji we wzorze (2) przez współczynnik τ Kendalla lub ρ_S Spearmana.

Liczba przekroczeń w okresie testowym oraz liczba przekroczeń statystyki Kupca we wszystkich analizowanych okresach dla trzech metod szacowania VaR jest zaprezentowana w tabelach 2 i 3.

Tabela 2

Liczba portfeli z przekroczeniami VaR powyżej poziomu tolerancji $\alpha = 0,01$
w grupie 405 badanych portfeli za pomocą trzech wariantów metody kowariancji

Okres		24.05.2000-11.10.2002			30.11.2004 -18.04.2007			2.12.2005-24.04.2008		
liczebność szeregu		125	250	500	125	250	500	125	250	500
ρ	wg przekroczeń (%)	28,89	9,14	4,94	35,06	34,32	20,74	36,3	60,99	6,42
	wg stat. Kupca	14	3	90	4	109	230	0	178	143
τ	wg przekroczeń (%)	32,84	9,38	7,9	35,8	35,31	46,17	37,04	61,98	12,35
	wg stat. Kupca	15	3	92	3	114	231	0	183	146
ρ_S	wg przekroczeń (%)	27,16	8,64	7,41	34,07	34,32	45,19	35,56	60,25	11,85
	wg stat. Kupca	11	3	89	3	110	224	0	177	141

Źródło: opracowanie własne.

Tabela 3

Liczba portfeli z przekroczeniami VaR powyżej poziomu tolerancji $\alpha = 0,05$
w grupie 405 badanych portfeli za pomocą trzech wariantów metody kowariancji

Okres		24.05.2000-11.10.2002			30.11.2004 -18.04.2007			2.12.2005-24.04.2008		
liczebność szeregu		125	250	500	125	250	500	125	250	500
ρ	wg przekroczeń (%)	13,09	0	0	4,69	1,23	0	0,49	4,94	0
	wg stat. Kupca	52	58	85	3	44	218	20	124	209
τ	wg przekroczeń (%)	13,09	0	0	0,49	1,73	0	4,69	5,68	0
	wg stat. Kupca	53	119	124	20	129	319	3	199	234
ρ_S	wg przekroczeń (%)	11,85	0	0	4,44	0,99	0	0,49	4,69	0
	wg stat. Kupca	59	135	122	4	130	316	26	193	232

Źródło: opracowanie własne.

Dokonując porównania otrzymanych przekroczeń można stwierdzić, że dla poziomu tolerancji $\alpha = 0,01$ najefektywniejsza okazała się metoda ze współczynnikiem ρ_S . W stosunku do klasycznej metody kowariancji otrzymano lepsze wyniki, zwłaszcza w krótszych próbach 125 i 250 obserwacji. Zastosowanie współczynnika τ wygenerowało wyniki nieznacznie gorsze niż dla metod ze współczynnikami ρ i ρ_S . Przy poziomie tolerancji $\alpha = 0,05$ metoda ze współczynnikiem ρ_S zawsze generowała mniej przekroczeń ponad poziom tolerancji niż metoda klasyczna. Zastosowanie współczynnika τ dało lepszy wynik dla okresu o trendzie rosnącym i 125 obserwacjach. W pozostałych przypadkach wyniki były porównywalne lub gorsze od rezultatów metody kowariancji ze współczynnikiem korelacji.

Na podstawie porównania wartości statystyki Kupca odnotowano prawie we wszystkich przypadkach ($\alpha = 0,05$) nieznaczny wzrost liczby portfeli z wartościami statystyki Kupca w obszarze krytycznym w stosunku do klasycznej metody kowariancji. Ponieważ znaczna część otrzymanych wyników nie zawierała żadnych przekroczeń ponad poziom tolerancji przy jednoczesnym występowaniu portfeli ze zbyt wysoką statystyką Kupca, można postawić hipotezę, iż dla poziomu tolerancji $\alpha = 0,05$ wyniki trzech metod są zbliżone. Różnice między liczbą portfeli z przekroczeniami, a liczbą portfeli wynikającą z testu Kupca można wyjaśnić otrzymywaniem liczby przekroczeń oscylujących wokół przyjętego poziomu tolerancji.

Druga część badania bezpośrednio poświęcona była wykorzystaniu kopuli w szacowaniu wartości zagrożonej.

W pierwszym kroku dokonano analizy wykorzystania kopuli do szacowania wartości zagrożonej dla dwóch wybranych portfeli: Impexmetal i PGF (z wagami odpowiednio 0,6 i 0,4) oraz Impexmetal i Rafako (z wagami 0,7 i 0,3). Szeregi obserwacji dla obu portfeli pochodziły z pierwszego okresu, tj. 24.05.2000-11.10.2002 i miały po 125 obserwacji. VaR szacowany był w przypadku pierwszego portfela przy poziomie ufności $\alpha = 0,01$ drugim zaś $\alpha = 0,05$. Wybrane portfele charakteryzowały się maksymalną liczbą przekroczeń w metodzie kowariancji spośród wszystkich badanych portfeli we wszystkich okresach (6 przekroczeń dla pierwszego portfela, 11 dla drugiego).

Do wyznaczenia oszacowania VaR zastosowano kopule Gaussa, Clayтона i Gumbel-Hougaarda. Symulacje przeprowadzono w dwóch wersjach. W pierwszej generowano 10000 razy 125-elementowe próbki o zadanej kopule, a następnie odczytywano poziom VaR. Ostateczna ocena była średnią otrzymanych wyników. W drugiej generowano jedną próbę o liczebności 10000 i odczytywano wartość zagrożoną.

Dla kopuli Gaussa procedura symulowania dziennych zmian wartości portfela w celu wyznaczenia VaR jest identyczna jak dobrze znana z literatury przedmiotu metoda symulacji Monte Carlo dla rozkładu normalnego.

Otrzymane wyniki analizy zamieszczone są w tabelach 4 i 5.

Tabela 4

Symulacja MC (próba o liczebności 10000)

kopuła	Impexmetal – PGF		Impexmetal – Rafako	
	liczba przekroczeń	statystyka Kupca	liczba przekroczeń	statystyka Kupca
Gaussa	6	11,758	11	5,733249
Gumbel-Hougaard	5	8,258217	8	1,615808
Clayton	2	0,7827239	8	1,615808

Źródło: opracowanie własne.

Tabela 5

Symulacja MC (próba o liczebności 125, 10000 iteracji dla kopuli Gaussa i 2000 dla pozostałych)

kopuła	Impexmetal – PGF		Impexmetal – Rafako	
	liczba przekroczeń	statystyka Kupca	liczba przekroczeń	statystyka Kupca
Gaussa	6	11,758	11	5,733249
Gumbel-Hougaard	6	11,758	10	4,130844
Clayton	3	2,632353	9	2,75

Źródło: opracowanie własne.

Estymacja wartości VaR za pomocą kopuli Gaussa dała takie same liczby przekroczeń jak metoda kowariancji. Jednakże metody z kopułą Gumbel-Hougaard i Claytona wygenerowały lepsze wyniki. Znaczącą redukcję liczby przekroczeń uzyskano dla obu portfeli wykorzystując kopułę Claytona.

Otrzymane wyniki w tych dwóch przypadkach dla kopuli Claytona okazały się obiecujące. Dlatego w drugim etapie badania dokonano estymacji wartości VaR za pomocą symulacji z wykorzystaniem kopuli Claytona dla szeregu 125 obserwacji z okresu 24.05.2000 – 11.10.2002 i obu poziomów tolerancji: $\alpha = 0,01$ oraz $\alpha = 0,05$. Wyniki tej analizy zawarte są w tabelach 6 oraz 7.

Tabela 6

Empiryczny rozkład procentowy przekroczeń w próbie testowej stu dni dla 405 badanych portfeli przy poziomie tolerancji $\alpha = 0.01$

liczba przekroczeń	metoda kowariancji (%)	kopula Claytona (%)
0	20,25	40,00
1	51,11	47,40
2	18,02	10,62
3	7,16	1,98
4	2,72	---
5	0,49	---
6	0,25	---

Źródło: opracowanie własne.

Tabela 7

Empiryczny rozkład procentowy przekroczeń w próbie testowej stu dni dla 405 badanych portfeli przy poziomie tolerancji $\alpha = 0.05$

liczba przekroczeń	metoda kowariancji (%)	kopula Claytona (%)
0	0,99	0,00
1	10,86	2,47
2	9,14	7,90
3	23,46	16,05
4	15,80	15,06
5	14,32	17,04
6	12,35	13,58
7	4,94	20,25
8	3,21	6,67
9	2,96	0,99
10	1,23	---
11	0,74	---

Źródło: opracowanie własne.

Przeprowadzone symulacje wykazują, że zaproponowana przez autora procedura szacowania wartości zagrożonej z wykorzystaniem kopuli Claytona, wygenerowała lepsze wyniki niż metoda kowariancji. Po pierwsze, dla obu poziomów tolerancji zaob-

serwowano, że symulacja za pomocą kopuli Claytona nie generuje wysokich wartości przekroczeń, tj. dla $\alpha = 0,05$ nie wystąpiło 10 i 11 przekroczeń, które pojawiły się w 1,97% przypadków w metodzie kowariancji, dla $\alpha = 0,01$ natomiast nie wystąpiło 4, 5 i 6 przekroczeń, które były w 3,46% przypadków z metody kowariancji. Po drugie, dla poziomu tolerancji $\alpha = 0,01$ zaobserwowano istotne zmniejszenie liczby przekroczeń, zaś wzrost liczby portfeli bez przekroczeń w porównaniu z wynikami metody kowariancji. Jednakże takiej prawidłowości nie ma przy poziomie $\alpha = 0,05$. W tym przypadku niepokojący jest wzrost liczby portfeli, dla których wystąpiło 7 przekroczeń. Jest to wartość, która wskazuje na przekroczenie ponad zakładany poziom. Warto jednak podkreślić, że jest to małe przekroczenie poziomu tolerancji przy jednoczesnym braku większych przekroczeń, tj. 10 i 11.

Na podstawie wartości statystyki Kupca w obu przypadkach metody dały podobne rezultaty. Dla poziomu $\alpha = 0,01$ nie zaobserwowano wartości statystyki w obszarze krytycznym dla metody z kopulą Claytona, dwie zaś dla metody kowariancji. Dla poziomu $\alpha = 0,05$ było to odpowiednio 10 oraz 11 wartości.

5. PODSUMOWANIE

Szacowanie wartości zagrożonej za pomocą kopuli, a także miar τ Kendalla i ρ_S Spearmana w przeprowadzonych symulacjach wskazuje, że otrzymane oszacowania są generalnie lepsze niż w metodzie kowariancji, w której zależność zmiennych zysków i strat składników portfela jest mierzona liniowym współczynnikiem korelacji. Oczywiście, opisane w artykule analizy stanowią wstęp do dalszych, pogłębionych badań. Należałoby zaproponowaną przez autora metodę szacowania VaR z wykorzystaniem kopuli Claytona przeanalizować dla innych typów szeregów danych oraz większej liczby danych. Przeprowadzone bowiem badanie z wykorzystaniem kopuli jest jedynie analizą dwóch przypadków. Otrzymane wyniki są obiecujące, lecz trudno na tym etapie badania formułować ogólne wnioski o skuteczności metod. Dalsze analizy dla większej liczby portfeli są kontynuowane przez autora. Niemniej jednak można postawić bardzo wiarygodną hipotezę, że zaproponowane metody są skuteczniejsze od metody kowariancji przy poziomie tolerancji $\alpha = 0,01$

Szkoła Główna Handlowa w Warszawie

LITERATURA

- [1] Bałamut T., [2002], *Metody estymacji VaR*, Materiały i studia, zeszyt nr 147, NBP, Warszawa.
- [2] Cherubini U., Luciano E., Vecchiato W., [2004], *Copula Methods in Finance*, John Wiley & Sons Ltd., West Sussex.
- [3] Jorion P., [2000], *Value at Risk. The new benchmark for managing financial risk*, McGraw-Hill, New York.
- [4] Kpanzou T.A., [2007], *Copulas in statistics*, African Institute for Mathematical Sciences, University of Stellenbosch.
- [5] Papla D., Piontek K., [2003], *Zastosowanie rozkładów α -stabilnych i funkcji powiązań (copula) w obliczaniu wartości zagrożonej (VaR)*, Akademia Ekonomiczna we Wrocławiu, Katedra Inwestycji Finansowych i Ubezpieczeń (http://www.kpiontek.ae.wroc.pl/DP_KP_VaR_copula.pdf).

- [6] Stryjek A., [2008], *Zastosowanie kopuli do szacowania wartości zagrożonej*, Badanie statutowe nr 03/S/0015/07, Szkoła Główna Handlowa, Warszawa.
- [7] Stryjek A., [2008], *Efekt dywersyfikacji ryzyka a efektywność klasycznych metod szacowania VaR dla portfeli GPW*, Badanie statutowe nr 03/S/0014/07, Szkoła Główna Handlowa, Warszawa.
- [8] Stryjek A., [2007], *Value at Risk a dywersyfikacja ryzyka*, Badanie statutowe nr 03/S/0060/06, Szkoła Główna Handlowa, Warszawa.
- [9] *Zarządzanie ryzykiem*, [2007], red. K. Jajuga, Wydawnictwo Naukowe PWN, Warszawa.

Praca wpłynęła do redakcji w maju 2009 r.

ZASTOSOWANIE MIAR ZALEŻNOŚCI ZMIENNYCH LOSOWYCH ORAZ KOPULI CLAYTONA I GUMBEL-HOUGAARDA DO SZACOWANIA WARTOŚCI ZAGROŻONEJ

Streszczenie

Opracowanie prezentuje możliwości, jakie daje kopula do szacowania wartości zagrożonej VaR. Autor przedstawia wyniki badania empirycznego dla portfeli spółek GPW w Warszawie. W badaniu dokonano porównania efektywności klasycznej metody kowariancji z dobrze znanymi z literatury przedmiotu oraz nowymi, zaproponowanymi przez autora metodami wykorzystującymi kopule Claytona i Gumbel-Hougaard.

Słowa kluczowe: zarządzanie ryzykiem, wartość zagrożona, kopula, symulacje Monte Carlo

APPLICATION OF RANDOM VARIABLES DEPENDENCE MEASURES AND CLAYTON AND GUMBEL-HOUGAARD COPULAS FOR ESTIMATING VALUE AT RISK

Summary

This paper shows the opportunities of copula for estimating Value at Risk (VaR). The author presents results of empirical research carried out for portfolios of stocks from Warsaw Stock Exchange. Efficiency of classical covariance method was compared with other well known in the literature and also new methods proposed by author using Clayton and Gumbel-Hougaard copulas.

Key words: risk management, Value at Risk, copula, Monte Carlo simulations

MICHAŁ KOLUPA, JOANNA PLEBANIĄK

O ELIMINACJI BŁĘDÓW W PROGNOZACH WYKONYWANYCH NA PODSTAWIE MODELU LEONTIEWA W PRZYPADKU AGREGACJI NIEDOSKONAŁEJ W SENSIE HATANAKI

Celem artykułu jest przedstawienie eliminacji błędów w prognozach wykonanych na podstawie modelu Leontiewa w przypadku, kiedy agregacja nie jest doskonała w sensie Hatanaki. Proponowane podejście wykorzystuje macierze brzegowe i zapewnia uzyskanie prognoz nie obciążonych błędami wynikającymi z niedoskonałej agregacji.

Dany jest schemat pierwotny Leontiewa, czyli trójka $(\mathbf{P}_X, \mathbf{X}, \mathbf{x})$, gdzie $\mathbf{P}_X = [x_{ij}]$ jest macierzą kwadratową stopnia n o elementach stanowiących przepływy międzygałęziowe z i -tej do j -tej gałęzi. Wektory $\mathbf{X}$ oraz $\mathbf{x}$ są n -wymiarowymi wektorami kolumnowymi o składowych odpowiednich równych X_i oraz x_i będących wartościami produkcji globalnej oraz końcowej i -tej gałęzi. Jest przy tym:

$$X_i = x_{i1} + x_{i2} + \dots + x_{in} + x_i, \quad i = 1, 2, \dots, n. \quad (1)$$

Oznaczmy symbolem $\mathcal{X}$ macierz diagonalną stopnia n taką, że:

$$\mathcal{X} = \text{diag}\{X_1, X_2, \dots, X_n\}. \quad (2)$$

Definiujemy macierz $\mathbf{A} = [a_{ij}]$ o wymiarach $n \times n$ w następujący sposób:

$$\mathbf{A} = \mathbf{P}_X \mathcal{X}^{-1}. \quad (3)$$

Zauważmy, że:

$$\begin{aligned} a_{ij} &\geq 0, \\ \sum_{j=1}^n a_{ij} &< 1, \quad \text{dla } i = 1, 2, \dots, n. \end{aligned} \quad (4)$$

Z (3) wynika, że:

$$\mathbf{P}_X = \mathbf{A} \mathcal{X}, \quad (5)$$

co po podstawieniu do (1) prowadzi do modelu Leontiewa. Ma on postać:

$$(\mathbf{I} - \mathbf{A}) \mathbf{X} = \mathbf{x}. \quad (6)$$

Na podstawie modelu Leontiewa danego wzorem (6), który będziemy nazywali pierwotnym, możliwe są do wykonania prognozy.

Prognoza prosta pierwszego rodzaju polega na wyznaczeniu wektora wartości przyrostu produkcji końcowej – $\Delta \mathbf{x}$, przy założeniu, że znamy macierz $\mathbf{I} - \mathbf{A}$ oraz wektor przyrostu wartości produkcji globalnej $\Delta \mathbf{X}$.

Z kolei prognoza drugiego rodzaju polega na wyznaczeniu wektora przyrostu wartości produkcji globalnej $\Delta \mathbf{X}$, kiedy znamy macierz $\mathbf{I} - \mathbf{A}$ oraz wektor przyrostu wartości produkcji końcowej $\Delta \mathbf{x}$.

Dodajmy jeszcze, że stopień n macierzy $\mathbf{I} - \mathbf{A}$ nazywamy stopniem pierwotnego modelu Leontiewa, który dany jest wzorem (6). Z kolei model postaci:

$$(\mathbf{I} - \mathbf{B}) \Delta \mathbf{Y} = \Delta \mathbf{y} \quad (7)$$

jest zagregowanym modelem Leontiewa, gdzie macierz $\mathbf{B}$ jest odpowiednikiem macierzy $\mathbf{A}$ podobnie jak r -wymiarowe wektory kolumnowe o składowych odpowiednio równych ΔY_i oraz Δy_i , $i = 1, 2, \dots, r$ stanowiące zagregowane wektory odpowiednio przyrostu wartości produkcji globalnej i końcowej.

Ponieważ stopień pierwotnego modelu Leontiewa jest wyższy od stopnia modelu wtórnego (zagregowanego) ($r < n$), przeto agregacja jest zabiegiem pożądanym. Niesie ona jednak pewne problemy, które niżej będą przedmiotem analiz.

Zatem analogicznie do schematu pierwotnego danego wzorem (1) mamy do czynienia z zagregowanym schematem Leontiewa:

$$Y_j = y_{j1} + y_{j2} + \dots + y_{jr} + y_j, \quad j = 1, 2, \dots, r, \quad (8)$$

zaś powyżej podane wielkości są zagregowanymi wielkościami stanowiącymi analogiczne wielkości do danych wzorami (2), (3), (4) i (5), czyli:

$$\mathbf{Y} = \text{diag}\{Y_1, Y_2, \dots, Y_r\}, \quad (9)$$

zaś macierz o wymiarach $r \times r$:

$$\mathbf{B} = \mathbf{P}_y \mathbf{Y}^{-1}. \quad (10)$$

ma elementy równe b_{ij} spełniające warunek:

$$b_{ij} \geq 0, \quad \sum_{j=1}^r b_{ij} < 1, \quad \text{dla } i = 1, 2, \dots, r. \quad (11)$$

Jeżeli model pierwotny ma stopień n , zaś model zagregowany jest stopnia r ($r < n$) ([3]), to schemat agregacji opisywany jest przez specjalną macierz $\mathbf{S}$ o wymiarach $r \times n$.

Odnotujemy jeszcze, że schemat agregacji stanowi zbiór reguł, według których łączy się gałęzie modelu pierwotnego w odpowiednie gałęzie modelu zagregowanego. Schemat taki nazywamy zgodnym, jeżeli spełnia następujące warunki [3]:

- a) każda gałąź modelu pierwotnego jest zaliczona do pewnej gałęzi modelu zagregowanego,
- b) jedna i ta sama gałąź modelu pierwotnego nie może być zaliczona do dwóch różnych gałęzi modelu zagregowanego,
- c) gałęzie zagregowane mają sens ekonomiczny.

Przypuśćmy, że pierwsze k_1 gałęzi modelu pierwotnego zostały zaliczone do pierwszej gałęzi modelu zagregowanego. Wówczas pierwszy wiersz macierzy $\mathbf{S}$ ma na węzłach $(1, 1), (1, 2), \dots, (1, k_1)$ jedynki, zaś pozostałe elementy tego wiersza są zerami.

Z kolei drugą gałąź modelu zagregowanego tworzą gałęzie o numerach $k_1 + 1, \dots, k_2$ modelu pierwotnego. Wówczas drugi wiersz macierzy $\mathbf{S}$ ma jedynki na miejscach $(2, k_1 + 1), \dots, (2, k_2)$, zaś pozostałe elementy tego wiersza są zerami, itd. i na koniec gałęzie o numerach $k_{r-1} + 1, \dots, k_r$ tworzą r -tą (ostatnią) gałąź w modelu zagregowanym. Wówczas elementy wiersza r -tego macierzy $\mathbf{S}$ stojące w węzłach $(r, k_{r-1} + 1), \dots, (r, k_r)$ są jedynkami, zaś pozostałe elementy tego wiersza są zerami.

Przypomnijmy jeszcze praktycznie ważne własności macierzy $\mathbf{S}$ ([4]):

1^0 rząd $\mathbf{S} = r$,

2^0 iloczyn $\mathbf{S}^T \mathbf{S}$ jest macierzą diagonalną,

3^0 jeżeli d_i oznacza liczbę jedynek występującą w r -tym wierszu macierzy $\mathbf{S}$, to i -ty diagonalny element tej macierzy jest równy d_i .

Obecnie zdefiniujemy agregację doskonałą.

W przypadku prognozy prostej pierwszego rodzaju agregacja $\mathbf{S}$ jest doskonała, jeżeli z warunku:

$$\mathbf{S} \Delta \mathbf{X} = \Delta \mathbf{Y} \quad (13)$$

spełnionego dla dowolnego wektora $\Delta \mathbf{X}$ wynika, że:

$$\mathbf{S} \Delta \mathbf{x} = \Delta \mathbf{y}. \quad (14)$$

Z kolei w przypadku prognozy prostej drugiego rodzaju agregacja $\mathbf{S}$ jest doskonała, jeżeli z warunku:

$$\mathbf{S} \Delta \mathbf{x} = \Delta \mathbf{y} \quad (15)$$

spełnionego dla dowolnego wektora $\Delta \mathbf{x}$ wynika, że:

$$\mathbf{S} \Delta \mathbf{X} = \Delta \mathbf{Y}. \quad (16)$$

Z kolei niżej podane twierdzenia dotyczą agregacji doskonałej.

Twierdzenie 1 (Hatanaka [2])

Warunkiem koniecznym i dostatecznym na to, aby agregacja $\mathbf{S}$ była doskonała dla prognozy prostej tak pierwszego jak i drugiego rodzaju jest spełnienie warunku:

$$\mathbf{SA} = \mathbf{BS}. \quad (17)$$

Posługiwanie się tym wzorem wymaga znajomości schematu pierwotnego. W praktyce wymóg ten nie zawsze jest spełniony.

Wówczas chcąc rozstrzygnąć, czy dana agregacja $\mathbf{S}$ jest doskonała korzystamy z twierdzenia Ary ([1]) podanego w wersji zamieszczonej w pracy [4].

Twierdzenie 2 (Ara)

Warunkiem koniecznym i wystarczającym na to, aby agregacja $\mathbf{S}$ była doskonała jest spełnienie równania:

$$\mathbf{SA} = \mathbf{SAS}^T(\mathbf{SS})^{-1}\mathbf{S}. \quad (18)$$

Posługiwanie się tym wzorem jest ułatwione z uwagi na poprzednio omówione własności macierzy $\mathbf{S}$.

Zakładamy, że agregacja $\mathbf{S}$ jest niedoskonała.

W przypadku prognozy prostej pierwszego rodzaju z warunku $\mathbf{S}\Delta\mathbf{X} = \Delta\mathbf{Y}$ spełnionego dla dowolnego wektora $\Delta\mathbf{X}$ wynika nie warunek (14), lecz warunek:

$$\mathbf{S}\Delta\mathbf{x} = \Delta\mathbf{y} + \delta, \quad (19)$$

gdzie δ jest niezerowym wektorem błędu.

W przypadku prognozy prostej drugiego rodzaju z warunku (15) spełnionego dla dowolnego wektora $\Delta\mathbf{x}$ wynika nie warunek (16), lecz warunek:

$$\mathbf{S}\Delta\mathbf{X} = \Delta\mathbf{Y} + \eta, \quad (20)$$

gdzie η jest niezerowym wektorem błędu.

Wektor δ może być wyznaczony na podstawie wzoru Czechowskiego (patrz [3]) postaci:

$$\delta = (\mathbf{BS} - \mathbf{SA}) \Delta\mathbf{X}. \quad (21)$$

Możliwe jest również wykorzystanie do tego celu macierzy brzegowej. I tak ma ona postać:

$$\mathbf{Q} = \begin{bmatrix} \mathbf{I}_n & \Delta\mathbf{X} \\ -(\mathbf{BS} - \mathbf{SA}) & \mathbf{0} \end{bmatrix}. \quad (22)$$

Po wykonaniu na macierzy $\mathbf{Q}$ przekształceń elementarnych typu β (macierz $-\mathbf{BS} - \mathbf{SA}$) przechodzi w macierz zerową) otrzymujemy:

$$\mathbf{Q} \sim \begin{bmatrix} (\mathbf{I}_n)^* & (\Delta\mathbf{X})^* \\ \mathbf{0} & \delta \end{bmatrix}, \quad (23)$$

co pozwala odczytać wektor δ .

W pracy [4] rozpatrzono przypadek kiedy $\delta = 0$, zaś macierz podstawowa układu (21) jest macierzą rzędu t ($0 \leq t \leq r$). Wówczas otrzymujemy wektor ΔX spełniający warunek:

$$(\mathbf{BS} - \mathbf{SA}) \Delta X = 0, \quad (24)$$

który daje wektor $\delta = 0$ mimo że $\mathbf{BS} \neq \mathbf{SA}$.

Z kolei w przypadku prognozy drugiego rodzaju wektor η spełnia warunek:

$$(\mathbf{I} - \mathbf{B}) \eta = (\mathbf{SA} - \mathbf{BS}) \Delta X. \quad (25)$$

Zauważmy, że nie znamy wektora ΔX i dlatego bezpośrednie skorzystanie z tego wzoru nie jest możliwe. Staje się ono możliwe po wykorzystaniu odpowiednich macierzy brzegowych.

Podamy dwa takie sposoby podejścia. Pierwszy polega na tym, iż możemy wyznaczyć wektor:

$$(\mathbf{SA} - \mathbf{BS}) \Delta X = (\mathbf{SA} - \mathbf{BS}) (\mathbf{I} - \mathbf{A})^{-1} \Delta x \quad (26)$$

nie wyznaczając przy tym wektora ΔX , lecz od razu iloczyn macierzy $(\mathbf{SA} - \mathbf{BS}) \Delta X$, czyli wektor $(\mathbf{I} - \mathbf{A})^{-1} \Delta x$.

Podkreślamy ten fakt z naciskiem. Jest to analogia do rozwiązywania układu równań liniowych $Ax = b$, przy czym chcąc uzyskać rozwiązanie nie wyznaczamy osobno macierzy odwrotnej do macierzy A , a następnie nie mnożymy macierzy A^{-1} przez wektor b , lecz od razu wyznaczamy wektor $A^{-1}b$. Do tego celu można wykorzystać odpowiednią macierz brzegową.

A zatem, chcąc wyznaczyć wektor stanowiący prawą stronę układu (25) najpierw skorzystamy z macierzy brzegowej postaci:

$$U = \begin{bmatrix} \mathbf{I} - \mathbf{A} & \Delta x \\ -(\mathbf{SA} - \mathbf{BS}) & 0 \end{bmatrix}. \quad (27)$$

Po wykonaniu na macierzy U przekształceń elementarnych typu α (macierz $\mathbf{I} - \mathbf{A}$ przechodzi w górną macierz trójkątną z jedynkami na głównej przekątnej) oraz przekształceń elementarnych typu β (macierz $-(\mathbf{SA} - \mathbf{BS})$ przechodzi w macierz zerową) dostajemy:

$$U \sim \begin{bmatrix} (\mathbf{I} - \mathbf{A})^* & (\Delta x)^* \\ 0 & (\mathbf{SA} - \mathbf{BS}) \Delta X \end{bmatrix}. \quad (28)$$

Stąd odczytujemy wektor $(\mathbf{SA} - \mathbf{BS}) \Delta X$ nie znając *a priori* wektora ΔX .

Znając już wektor $(\mathbf{SA} - \mathbf{BS}) \Delta X$ rozwiązujemy układ (25). Do tego celu możemy wykorzystać macierz brzegową postaci:

$$\mathbf{V} = \begin{bmatrix} \mathbf{I} - \mathbf{B} & (\mathbf{S}\mathbf{A} - \mathbf{B}\mathbf{S}) \Delta \mathbf{X} \\ -\mathbf{I}_r & \mathbf{0} \end{bmatrix}. \quad (29)$$

Po wykonaniu na macierzy $\mathbf{V}$ przekształceń elementarnych typu α (macierz $\mathbf{I} - \mathbf{B}$ przechodzi w górną macierz trójkątną z jedynkami na głównej przekątnej) oraz przekształceń elementarnych typu β (macierz $-\mathbf{I}_r$ przechodzi w macierz zerową) dostajemy:

$$\mathbf{V} \sim \begin{bmatrix} (\mathbf{I} - \mathbf{B})^* & [(\mathbf{S}\mathbf{A} - \mathbf{B}\mathbf{S}) \Delta \mathbf{X}]^* \\ \mathbf{0} & \boldsymbol{\eta} \end{bmatrix}. \quad (30)$$

Stąd odczytujemy wektor $\boldsymbol{\eta}$ będący rozwiązaniem układu (25) (przypomnijmy nie była w tym przypadku konieczna *a priori* znajomość wektora $\Delta \mathbf{X}$).

Zauważmy, że znacznie prościej jest wykorzystać macierz brzegową $\mathbf{F}$ postaci ([5]):

$$\mathbf{F} = \begin{bmatrix} \mathbf{I} - \mathbf{A} & \Delta \mathbf{x} \\ -\mathbf{S} & \mathbf{0} \end{bmatrix} \quad (31)$$

i wykonać na niej przekształcenia elementarne typu α (macierz $\mathbf{I} - \mathbf{A}$ przechodzi w górną macierz trójkątną z jedynkami na głównej przekątnej) oraz przekształceń elementarnych typu β (macierz $-\mathbf{S}$ przechodzi w macierz zerową). Wówczas:

$$\mathbf{F} \sim \begin{bmatrix} (\mathbf{I} - \mathbf{A})^* & (\Delta \mathbf{x})^* \\ \mathbf{0} & \Delta \mathbf{Y} \end{bmatrix}, \quad (32)$$

gdzie zawsze $\Delta \mathbf{Y} = \mathbf{S}\Delta \mathbf{X}$.

Politechnika Radomska
Szkoła Główna Handlowa

LITERATURA

- [1] Ara K., [1959], *The Aggregation Problem In Input – Output Analysis*, „Econometrica” 27.
- [2] Czechowski T., [1960], *Kilka uwag o agregacji w modelu Leontiewa*, „Przegląd Statystyczny” 3, Warszawa.
- [3] Hatanaka M., [1952], *Note on Consolidation with Leontief System*, „Econometrica” 4.
- [4] Kolupa M., [1969], *O eliminacji błędów w prognozach wykonywanych na podstawie modelu Leontiewa*, „Przegląd Statystyczny” 2, Warszawa.
- [5] Łukasik J. (obecnie Plebaniak), [1983], *Zastosowanie macierzy brzegowych do prognozowania na podstawie modelu Leontiewa*, Monografie i Opracowania SGPiS, Warszawa.

O ELIMINACJI BŁĘDÓW W PROGNOZACH WYKONYWANYCH NA PODSTAWIE
MODELU LEONTIEWA W PRZYPADKU AGREGACJI NIEDOSKONAŁEJ W SENSIE HATANAKI

Streszczenie

Dany jest pierwotny model Leontiewa stopnia n postaci:

$$(\mathbf{I} - \mathbf{A}) \Delta \mathbf{X} = \Delta \mathbf{x}.$$

Dokonujemy agregacji $\mathbf{S}$ (macierz o wymiarach $r \times n$ ($r < n$)) uzyskując wtórny (zagregowany) model Leontiewa stopnia r postaci:

$$(\mathbf{I} - \mathbf{B}) \Delta \mathbf{Y} = \Delta \mathbf{y}.$$

Jeżeli agregacja nie spełnia $\mathbf{BS} = \mathbf{SA}$, to w prognozach prostych zarówno pierwszego, jak i drugiego rodzaju pojawiają się błędy.

Dla prognozy pierwszego rodzaju wektor błędów δ możemy wyznaczyć z warunku Czechowskiego:

$$(\mathbf{BS} - \mathbf{SA}) \Delta \mathbf{X} = \delta.$$

Jeżeli przy tym rząd macierzy $\mathbf{BS} - \mathbf{SA}$ ($r(\mathbf{BS} - \mathbf{SA})$) jest równy t ($0 < t \leq r$), to istnieją wektory $\Delta \mathbf{X}$, dla których $\delta = \mathbf{0}$ mimo że agregacja nie jest doskonała ($\mathbf{BS} \neq \mathbf{SA}$).

Dla prognozy drugiego rodzaju konstruując macierz brzegową postaci:

$$\mathbf{F} = \begin{bmatrix} \mathbf{I} - \mathbf{A} & \Delta \mathbf{X} \\ -\mathbf{S} & \mathbf{0} \end{bmatrix}$$

zapewniamy, że zawsze $\mathbf{S} \Delta \mathbf{X} = \Delta \mathbf{Y}$ mimo że $\mathbf{BS} \neq \mathbf{SA}$.

W przypadku prognozy drugiego rodzaju wektor η ($\mathbf{BS} \neq \mathbf{SA}$ – agregacja nie jest doskonała) może być wyznaczony z równania:

$$(\mathbf{I} - \mathbf{B}) \eta = (\mathbf{SA} - \mathbf{BS}) \Delta \mathbf{X}.$$

Jako, że nie znamy *a priori* $\Delta \mathbf{X}$, możemy go wyznaczyć wykorzystując kolejno macierze brzegowe postaci:

$$\mathbf{U} = \begin{bmatrix} \mathbf{I} - \mathbf{A} & \Delta \mathbf{x} \\ -(\mathbf{SA} - \mathbf{BS}) & \mathbf{0} \end{bmatrix} \text{ oraz } \mathbf{V} = \begin{bmatrix} \mathbf{I} - \mathbf{B} & (\mathbf{SA} - \mathbf{BS}) \Delta \mathbf{X} \\ -\mathbf{I}_r & \mathbf{0} \end{bmatrix}$$

i wtedy wektor η może być wyznaczony.

Słowa kluczowe: model Leontiewa, prognozy ekonometryczne, agregacja

ON ELIMINATING ERRORS OF FORECASTS OBTAINED FROM THE LEONTIEF MODEL
FOR THE CASE OF IMPERFECT HATANAKI'S AGGREGATION

Summary

In the following article, there is presented a primary Leontiew model of order $n \times n$:

$$(\mathbf{I} - \mathbf{A}) \Delta \mathbf{X} = \Delta \mathbf{x}.$$

Aggregation **S** is performed (matrix **S** is a matrix of order $r \times n$ ($r < n$)) and the effect of it is an aggregated Leontiew model of rank r :

$$(\mathbf{I} - \mathbf{B}) \Delta \mathbf{Y} = \Delta \mathbf{y}.$$

If aggregation does not fulfil $\mathbf{BS} = \mathbf{SA}$, then in simple prognoses, both of the first and second kind, some errors appear.

For the prognosis of the first kind, the errors vector δ can be determined from Czechowski's equation:

$$(\mathbf{BS} - \mathbf{SA}) \Delta \mathbf{X} = \delta.$$

If the rank of matrix $\mathbf{BS} - \mathbf{SA}$ ($r(\mathbf{BS} - \mathbf{SA})$) equals t ($0 < t \leq r$), there are $\Delta \mathbf{X}$ vectors, for which $\delta = \mathbf{0}$ despite the fact that aggregation is not perfect i.e. $\mathbf{BS} \neq \mathbf{SA}$.

For the prognosis of the second kind a bordered matrix:

$$\mathbf{F} = \begin{bmatrix} \mathbf{I} - \mathbf{A} & \Delta \mathbf{X} \\ -\mathbf{S} & \mathbf{0} \end{bmatrix}$$

is constructed; then always $\mathbf{S}\Delta \mathbf{X} = \Delta \mathbf{Y}$ in spite of the fact that $\mathbf{BS} \neq \mathbf{SA}$.

In case of the prognosis of the second kind the η vector ($\mathbf{BS} \neq \mathbf{SA}$ – aggregation is not fulfilled) can be derived from the equation:

$$(\mathbf{I} - \mathbf{B}) \eta = (\mathbf{SA} - \mathbf{BS}) \Delta \mathbf{X}.$$

As $\Delta \mathbf{X}$ is *a priori* not known, the following bordered matrices are applied in sequence:

$$\mathbf{U} = \begin{bmatrix} \mathbf{I} - \mathbf{A} & \Delta \mathbf{x} \\ -(\mathbf{SA} - \mathbf{BS}) & \mathbf{0} \end{bmatrix} \text{ oraz } \mathbf{V} = \begin{bmatrix} \mathbf{I} - \mathbf{B} & (\mathbf{SA} - \mathbf{BS}) \Delta \mathbf{X} \\ -\mathbf{I}_r & \mathbf{0} \end{bmatrix}$$

and the η vector can thus be determined.

Key words: Leontief model, econometric forecasts, aggregation.

ANTONI SMOLUK

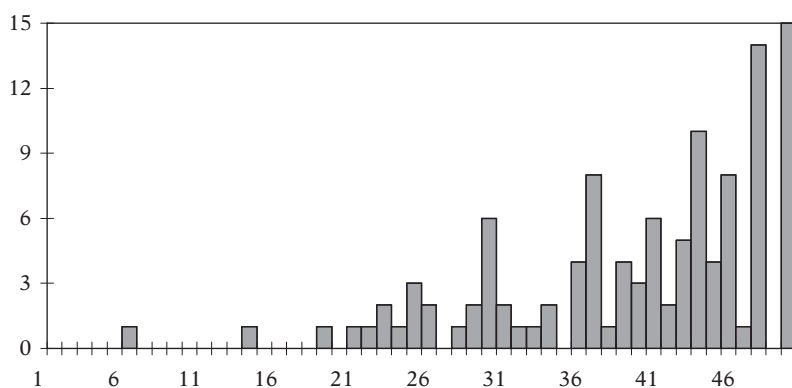
40 LAT DZIAŁALNOŚCI NAUKOWEJ PROFESORA JÓZEFA HOZERA, CZYLI O METROLOGII EKONOMICZNEJ

Gratia gratiam parit

1. Obecnego profesora zwyczajnego Józefa Hozera poznałem około 40 lat temu na konferencjach organizowanych przez Zbigniewa Pawłowskiego; był wtedy młodziutkim skromnym chłopcem, chyba bezpośrednio po studiach. Promieniowała od niego życzliwość, dobroć i wrodzona kultura. Ta przelotna znajomość może byłaby bez znaczenia, gdyby nie późniejsze, niemal systematyczne spotkania na konferencjach organizowanych w okolicy Świnoujścia przez Uniwersytet Szczeciński, a wcześniej Politechnikę. Spotkania te, raczej w wąskim gronie, przenikała atmosfera zrozumienia i współpracy; nie było strachu przed zjadliwą krytyką, czy obaw przed popełnieniem błędu. Konferencje szczególnie stymulowały rozwój młodych pracowników oraz ułatwiały im życie w świecie nauki poprzez prezentację szerokiego spektrum tematów badawczych; były swego rodzaju sceną pierwszych prób w nauce. Dobry przykład zbiorowego wysiłku we wspólnie prowadzonych badaniach. W wyniku tego systematycznego i systemowego działania ośrodek szczeciński wyrastał na wiodące centrum mikroekonometrii. Trzeba zaznaczyć, że cieszy się on – profesor Hozer – wielkim autorytetem i budzi szacunek nie tylko młodych i niedoświadczonych kolegów, ale także jest mężem zaufania całego środowiska naukowego, polityków, administracji lokalnej oraz centralnej. Można go śmiało nazwać wielkim mecenasem nauki, który wspiera naukę przez to, że potrafi sprzedać jej wyniki. *Dobrze chmielowi, gdy się trzyma tyki*. Spotykaliśmy się także niemal corocznie na zakopiańskich konferencjach organizowanych przez profesora Aleksandra Zeliasia oraz na wielu zebraniach Komitetu Ekonometrii i Statystyki Polskiej Akademii Nauk. Na jednym z tych komitetów wystąpił w obronie zaproponowanej problematyki badawczej, którą władze komitetu próbowały odsunąć na boczny tor – chodzi o pomiar w ogóle, a metrologię ekonomiczną w szczególności. Nie ma nauki bez pomiaru, a pomiar jest przecież niczym innym jak homeomorfizmem natury w strukturę formalną – system relacyjny; pomiar to o wiele więcej niż zwykły eksperyment – to hipotezy i teoria.

2. Ekonometria światowa i w ślad za nią ekonometria polska skupia się głównie na badaniu własności modeli. Szkoła szczecińska profesora Hozera traktuje modele inaczej – jako narzędzie poznania świata fizycznego, życia społecznego i gospodarczego.

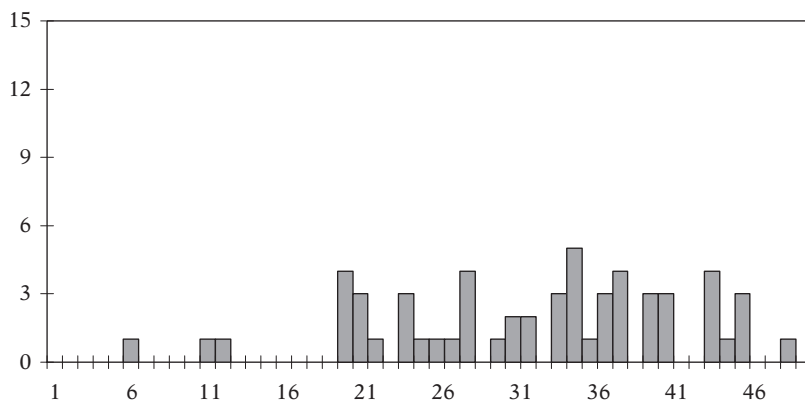
Specjalnością Szczecina jest mikroekonometria, czyli zastosowanie metod ekonometrycznych w przedsiębiorstwie. Auditing jest najprostszym zastosowaniem matematyki; prosty przegląd ksiąg rachunkowych może być i jest źródłem wiedzy o działalności firmy i jej kondycji finansowej. Nieuczciwych sprzedawców łatwo rozpoznać po sprawozdaniach handlowych.



Rysunek 1. Próba mało prawdopodobna

Źródło: A. Smoluk (red.), (2000).

Każdy wybór z populacji generalnej jest próbą losową. Próba zgodna ze strukturą populacji jest próbą reprezentacyjną. Próbą losową są również samochody, które w określonym dniu uległy wypadkowi. Populację generalną w tym przypadku stanowią wszystkie samochody, które w tym dniu były w ruchu. Próby mało prawdopodobne dostarczają niekiedy również wielu cennych informacji. Egzamin definiuje pewien preporządek – relację zwrotną i przechodnią – w badanej grupie. Egzamin jest pomiarem. Ten z kandydatów jest lepszy, który ma więcej punktów. Z populacji kandydatów na studentów pobrano dwie próby. Wyniki egzaminu z matematyki, w skali od 0 do 50 punktów, dla próby liczącej 113 osób, przedstawiono na rys. 1 (na osi poziomej oznaczono punkty, a na pionowej częstotliwości – liczbę kandydatów). Jeżeli przyjmiemy, że w populacji generalnej jest normalny rozkład wiadomości, to ta próba nie jest zgodna ze strukturą populacji; 15 kandydatów otrzymało maksymalną liczbę punktów. Albo kandydaci z tej komisji byli wyjątkowo zdolni, albo egzamin nie był starannie przeprowadzony (kandydaci chyba korzystali z jakiegoś wsparcia). Takie są wnioski z analizy rozkładu. Próba druga, licząca 57 kandydatów (rys. 2), wydaje się zgodna ze strukturą populacji. Moda w tej próbie równa się 35 punktom.



Rysunek 2. Próba typowa

Źródło: A. Smółuk (red.), (2000).

Nauka ma nawet metody weryfikacji rzetelności wyborów i prac komisji wyborczych. Dane empiryczne podlegają prawu Benforda sformułowanego przez amerykańskiego fizyka Franka Benforda w latach trzydziestych minionego wieku. Cyfry 1 i 2 są uprzywilejowane; od nich bowiem zaczynają się w przeważającej części wszystkie dane [1]. Jest to norma naukowa. Ekonometria bada prawidłowości życia gospodarczego, jej cechą charakterystyczną jest matematyka, ściślej, język formalny jako metoda; ekonometria jest specyficznie ujętą statystyką ekonomiczną. Ta specyfika sprowadza się do badania dynamiki procesów ekonomicznych. Początków ekonometrii jako nauki można dopatrzeć się w pierwszej ćwierci XX wieku. W 1926 r. Ragnar Frisch użył terminu ekonometria w rozumieniu dzisiejszym; nazwa ta jednak była znana dużo wcześniej i to użyto ją w języku polskim. Pewnie z tego powodu uszła uwadze opinii światowej, ponieważ polskiej literatury ekonomicznej świat nie studiuje. Paweł Ciompa (1867-1913) w pracy *Zarys ekonometrii i teorya buchalteryi* z 1910 r. wprowadził tę piękną nazwę, którą dziś wywodzi się od biometrii. Ekonometrię wykorzystywano przy prognozach koniunktury gospodarczej; stosowano układy wskaźników i indeksów różnego rodzaju zależnych naturalnie od gustu autora czy szkoły, z której się wywodził. Te parametry, szumnie zwane – by podbudować się prestiżem meteorologii – barometrami ekonomicznymi, były użyteczne, gdy trwała stabilna zrównoważona sytuacja, czyli gdy była dobra koniunktura. Kryzys lat 1929-1933 był zaskoczeniem dla wszystkich indeksów. Tak być musiało, gdy praw nauki nie znano, a stosowano w miejsce nich przetwarzanie danych, czyli głośny *data mining*. Rachunku korelacyjnego użył Galton do wyliczania wysokości dzieci w zależności od wzrostu rodziców. Zależność dwóch cech od siebie jest podstawą wielu praw przyrody. Każda funkcja ma charakter zależności korelacyjnej. Jest to jedna z najważniejszych metod biologii, antropologii i psychologii rozwinięta przez Pearsona, Spearmana i innych. Pomiar jest ważny, liczby niosą cenne informacje, ale bez logiki, bez znajomości porządku natury, z samych liczb mało wynika i to nawet wtedy, gdy

rachujemy na najlepszych komputerach. Trwałym osiągnięciem ekonometrii są metody wygładzania danych zaadoptowane głównie w pierwszym rzędzie z astronomii, później geodezji i biologii. Badanie szeregów czasowych jest zajęciem pożytecznym, by móc śledzić tendencje rozwojowe, jednakowoż na te badania nałożyło się ciężkie brzemie formalne, które zdominowało całą ekonometrię. Tworzymy metoda po metodzie nowe rodzaje rachunków, a nie szukamy relacji rządzących gospodarką, relacji użytecznych przy decyzjach organów administracyjnych i towarzystw gospodarczych. Formalizm zniszczy ekonometrię jako naukę; jest to nauka o modelach, czyli o niczym. Każda nauka ma substrat w świecie fizycznym; bez podłoża matematycznego nie ma jednak nauki, ale źródłem jest zawsze problem badawczy, a metoda jego pochodną; metoda jest w zadaniu, którego rozwiązania szukamy. Naturalne i optymalne rozwiązanie zadania jest zawsze częścią problemu, tkwi w zadaniu. Metodę tworzymy rozwiązując problem; ona sama się pojawia, gdy dobrze zrozumiemy zadanie, gdy pójdziemy drogą naturalną do celu.

3. Spośród znanych ekonometryków na wspomnienie zasługują przede wszystkim nazwiska: Frisch, Timbergen, Wold, Koopmans, Klein, Granger i jeszcze kilka innych. Profesor Hozer był na stażu naukowym u Kleina i jemu zaprezentował swoje odkrycie. Chodzi o *regulę pięciu*; jest to norma uniwersalna – ogólna prawidłowość obowiązująca na obecnym etapie rozwoju cywilizacyjnego. Na każde pięć gospodarstw domowych przypada jedno przedsiębiorstwo. Wielki Klein początkowo miał zastrzeżenia co do tej prawidłowości, ale po objaśnieniach i pokazaniu siły prognostycznej tego prawa, zmienił zdanie i chyba uwierzył w jego słuszość. Prawo Hozera ma podstawy psychologiczne i socjologiczne. W dużym zbiorowisku osób pozostawionych samym sobie tworzą się samorzutnie podgrupy; najczęściej są to grupki pięcioosobowe. Profesor Reinhard Selten – słynny noblista, doktor *honoris causa* Uniwersytetu Ekonomicznego we Wrocławiu – zauważył, że koalicje w grach są także najczęściej pięcioosobowe i rzadko przewyższają tę liczbę. Liczba 5 jest uprzywilejowana w popularnym zwrocie: *do policzenia wystarczą palce jednej ręki*. Mówimy tak wtedy, gdy mnogość przedmiotów jest nieliczna, gdy ich liczbę poznajemy jednym rzutem oczu. W przyrodzie istnieją pewne stałe uniwersalne, takie jak liczba Archimedesesa π będąca stosunkiem obwodu koła do jego średnicy, liczba Eulera e będąca kapitałem, jaki otrzymamy po roku z banku stosującego kapitalizację ciągłą, gdy włożyliśmy jednostkę przy procencie 100, stała Plancka h – kwant działania równy w przybliżeniu $2\pi \times 10^{-27} J \times s$ i wiele innych. Stałą fizyczną, powszechnie znaną, jest prędkość światła w próżni, a także wielkość ładunku elementarnego. Profesor Hozer słusznie uważa, że w ekonomii muszą być tak samo wielkości stałe, uniwersalne parametry, normy, standardy. Należy ich szukać. Taką wielkością jest właśnie liczba 5, ale taką wielkością może być również 8% lansowane przez profesora Mieczysława Dobiję jako naturalne oprocentowanie kredytu. Te osiem procent profesora Dobiji może mieć głębsze uzasadnienie w zjawisku Gibssa (*Gibbs phenomenon*). Jeżeli okresowa funkcja jest regularna, w uproszczeniu oznacza to, że jest w jednym punkcie nieciągła, a poza tym punktem jest gładka, w punkcie nieciągłości ma granice obustronne i jej wartość równa się średniej arytmetycznej granic lewostronnej i prawostronnej, to szereg Fouriera takiej funkcji jest (po wyłączeniu punktu nieciągłości z pewnym otoczeniem) jednostajnie zbieżny do tej funkcji. Szereg ten jest

zbieżny do tej funkcji, w zwykłym sensie – w każdym punkcie, ale jego zachowanie w toczeniu punktu nieciągłości jest niezwykle. Wartości szeregu Fouriera w punkcie nieciągłości są zbieżne do średniej arytmetycznej granicy lewostronnej i prawostronnej, ale by przeskoczyć rozstęp pomiędzy tymi granicami szereg bierze prawie dziewięcioprocentowy rozbieg z lewej strony i przeskakuje o taką samą część tego rozstępu powyżej potrzeby – do góry dla funkcji rosnącej. Kredyt jest nieciągłością rozwoju, stąd owa stała Gibbsa

$$\gamma = -\frac{1}{\pi} \int_{\pi}^{\infty} \frac{\sin x}{x} dx = 0.089\dots$$

może być argumentem za ośmioma procentami profesora Dobiji.

4. Teoria liczb wydaje się nauką bardzo odległą od ekonomii, jednakowoż relacja podzielności, podstawowe pojęcie tej teorii jest także sercem ekonomii. Podzielność jest abstrakcyjnym ujęciem relacji preferencji [5]. Zdaje sobie z tego sprawę profesor Hozer i swych pracowników zachęca do wszechstronnych studiów w tym także nad fizyką i teorią liczb. W jego katedrze ukazała się wartościowa pozycja poświęcona fundamentalnym i pięknym działom abstrakcyjnej teorii liczb. Ciekawy i trudny problem arytmetyczny *partitio numerorum* jest związany również z pojęciem normy i wzorca. Najpierw przypomnę, że partycją – podziałem – liczby naturalnej n różnej od zera – zero nie ma partycji – jest rosnący n -wyrazowy, a więc skończony ciąg liczb naturalnych, sumujących się do liczby n . Partycją liczby 5 jest wektor $(0, 0, 1, 1, 3)$, albowiem $1 + 1 + 3 = 5$. Partycje definiują pragmatykę systemu monetarnego. Wiadomo, że dobry pieniądz to silny pieniądz. Ale system jednostek pieniężnych powinien mieć także inną zaletę. Każdą kwotę powinno się w nim łatwo odliczyć w różnym układzie banknotów. Im więcej możliwości, tym lepiej, tym system monetarny z punktu widzenia pragmatycznego wygodniejszy. Wyznaczanie nominałów pieniężnych nie jest więc czynnością rutynową, ale tworzeniem przyjaznego systemu obiegowego. System pieniężny jest wygodny w użyciu, jeśli nominały grubsze dają się łatwo – i to różnorodnie – podzielić na mniejsze. Moneta 2 zł, pomijamy bilon o nominałach groszowych, rozmienna się tylko w jeden sposób $2 = 1 + 1$. Moneta 5 zł może być rozmienniona już na trzy sposoby: $(1, 1, 1, 1, 1)$, $(0, 1, 1, 1, 2)$, $(0, 0, 1, 2, 2)$, czyli na pięć monet jednozłotowych lub na trzy monety jednozłotowe i jedną dwuzłotową, albo wreszcie na jedną złotówkę i dwie dwuzłotówki. Innych podziałów pieniężnych nie ma, chociaż partycji liczby 5 jest więcej. Podział $(0, 0, 0, 0, 5)$ – bez podziału – będziemy uważać za trywialny. Banknot dziesięciozłotowy można rozmiennić na dziesięć sposobów nietrywialnych. Partycje nietrywialne są właśnie dopuszczalnymi w systemie pieniężnym podziałami. Ile jest partycji dopuszczalnych dowolnej liczby naturalnej? Na ile sposobów można opłacić rachunek stu złotych? Metodę, którą należy się posłużyć, daje teoria funkcji tworzących. Funkcje tworzące to szeregi potęgowe o współczynnikach całkowitych. Dla polskiego systemu pieniężnego funkcją tworzącą jest szereg potęgowy

$$f(x) = 1 + a_1x + a_2x^2 + \dots$$

będący iloczynem ośmiu szeregów

$$f_i(x) = 1 + x^i + x^{2i} + x^{3i} + \dots, \quad i \in \{1, 2, 5, 10, 20, 50, 100, 200\}.$$

Aby odpowiedzieć na pytanie, na ile sposobów można odliczyć kwotę 1000 zł, należy obliczyć współczynnik a_{1000} funkcji f . Tyle jest bowiem dopuszczalnych partycji liczby 1000 w polskim systemie pieniężnym. Kwotę 7 zł możemy zapłacić w kilku pulach: albo układem (1, 1, 1, 1, 1, 1, 1), czyli siedmioma monetami jednozłotowymi (tej partycji odpowiada wyraz x^7 funkcji tworzącej f_1), albo układem (0, 1, 1, 1, 1, 1, 2) (tej partycji odpowiada wyraz x^5 funkcji f_1 oraz wyraz x^2 funkcji f_2), albo układem (0, 0, 1, 1, 1, 2, 2) (tej partycji odpowiada wyraz x^3 funkcji f_1 oraz wyraz x^4 funkcji f_2), albo układem (0, 0, 0, 0, 1, 1, 5) (tej partycji odpowiada wyraz x^2 funkcji f_1 oraz wyraz x^5 funkcji f_5), albo układem (0, 0, 0, 1, 2, 2, 2) (tej partycji odpowiada wyraz x funkcji f_1 oraz wyraz x^6 funkcji f_2), albo wreszcie układem (0, 0, 0, 0, 0, 2, 5) (tej partycji odpowiada wyraz x^2 funkcji f_2 oraz wyraz x^5 szeregu f_5). Tak więc dopuszczalnych partycji liczby 7 jest 6, czyli na sześć sposobów możemy – w polskim systemie pieniężnym – zapłacić siedmioletowy rachunek; ciągle potrzeba pamiętać, że pomijamy monety groszowe: 1 gr, 2 gr, 5gr, 10 gr, 20 gr i 50 gr.

W praktycznym systemie pieniężnym monety i banknoty powinny być łatwo rozróżnialne i umożliwić wygodną regulację każdego rachunku. W aktualnie obowiązującym polskim systemie pieniężnym w parach: 2 gr i 5 gr, 20 gr i 50 gr oraz 2 zł i 5 zł monety są do siebie podobne. Łatwo więc o pomyłkę; z tego powodu nominałów nie może być dużo. Wydaje się, że ergonomiczny system pieniężny powinien utrzymać równowagę pomiędzy różnorodnością i rozróżnialnością znaków pieniężnych; wystarczy operować tylko czterema monetami (dla wygody posługujemy się tu polskimi jednostkami): 1 gr, 5 gr, 10 gr i 50 gr oraz pięcioma banknotami: 1 zł, 5 zł, 10 zł, 50 zł i 100 zł. W przypadku ogólnym, gdy dopuścimy dowolnie wysokie nominały, będzie to ciąg nieskończony postaci (1, 5, 10, 50, 100, 500, 1000, 5000, ...). Wyraz ogólny a_n tego ciągu może być określony indukcyjnie: $a_{2n} = 10^n$, $a_{2n+1} = 5a_{2n}$, $n \in N$. Wysokie nominały banknotów nie są, przy stabilnej walucie, wskazane. Duże rachunki reguluje się przelewami bankowymi. Liczba trzy jest także normą sprawnego ergonomicznego systemu pieniężnego; mamy bowiem nominały: 1, 5, 10; dalej ten układ jednostek mnożymy wielokrotnie przez dziesięć – w zależności od potrzeb rynku pieniężnego.

5. Profesor Józef Hozer jest uczniem inicjatora badań ekonometrycznych w Polsce profesora Zbigniewa Pawłowskiego. Prof. Z. Pawłowski miał wrodzoną kulturę, wyniósł z domu rodzinnego dobrą kindersztubę, był pod wpływem języka francuskiego i francuskiej cywilizacji. Nie było w nim wulgarności, był subtelny i delikatny. Język jego był piękny i prosty, nigdy nie popisywał się swoją wiedzą i uczonością. Dobierał podobnych do siebie uczniów i współpracowników; takim właśnie jest profesor Józef Hozer, takim jest profesor Oskar Starzeński i tacy są inni ulubieńcy Pawłowskiego. Od Pawłowskiego promieniowała na otoczenie szlachetność, zawsze starał się zrozumieć bliźnich, nigdy nie zniżył się do małostkowości. Profesor Pawłowski – ojciec polskiej ekonometrii – widział tę naukę jako dział ekonomii, specjalny dział ekonomii stosujący probabilistykę i matematykę. Dla Pawłowskiego celem badania była elastyczność dochodowa popytu,

funkcja produkcji czy metody optymalizacyjne z programowaniem liniowym na czele. Owszem w nauce ważna jest doskonałość i perfekcja, ale pierwszym powołaniem nauki jest poznanie rzeczywistości; poznanie rzeczywistości, a nie tworzenie scenariuszy i modeli, może nawet i pięknych, ale bezwartościowych, gdy nic praktycznego z nich nie wynika. Przy okazji zaznaczę, że rodzina profesora Hożera wywodzi się z Wiednia, a w dawnej Warszawie, jeszcze w dziewiętnastym wieku, był słynny ogrodnik tego nazwiska. W teorii zarządzania mówi się o zaufaniu przełożonego do podwładnych i odwrotnie – podwładnych do przełożonego. Profesor Hozer sprawuje swą funkcję przez zaufanie, którym darzy swych kolegów i współpracowników. Zaufanie jest bowiem podstawą każdej instytucji i zasadniczym czynnikiem ładu społecznego. Wiedział o tym lubiany cesarz Franciszek Józef, który był jaśnie oświeconym panem przodków profesora. Cesarz darzył pełnym zaufaniem swych urzędników. Jeden z jego generałów, nielubiany przez sztab generalny, budował twierdzę gdzieś daleko na Bałkanach. W sprawozdaniu finansowym napisał krótko, po spartańsku: 15 milionów guldenów dostałem i 15 milionów guldenów wydałem. Oficerowie ze sztabu nie byli zadowoleni lakonicznością tego raportu i poprosili o szczegóły; generał powtórzył poprzednie sprawozdanie i dodał: *A kto temu nie wierzy niech mnie pocałuje w...* Naturalnie o zdarzeniu doniesiono cesarzowi. – *Ja wierzę* – miał odpowiedzieć jaśnie pan. Profesor Hozer żartował z teorii *Übermenscha*, stworzonej przez Friedricha Nietzschego, który chciał być poza dobrem i złem. Dobra nie ma, zła także nie ma, jest tylko moja wola i to, co się mnie podoba. Takie stanowisko jest dziś powszechnie akceptowane, chociaż wszyscy oficjalnie uważają filozofię Nietzschego za wymysł szalonego geniusza. Czymże innym są bowiem parady równości, walka o legalność wszelkich badań biogenetycznych i demonstracje pań domagających się pełnej wolności w dysponowaniu swym ciałem jeśli nie woluntaryzmem? Akceptujemy to, co jest nam w danej chwili wygodne; nie ma sumienia, nie ma wstydu, jest skuteczność. Jeżeli osiągnąłem to, co chciałem, to prawda jest po mojej stronie. Przyjdzie jednak czas rozrachunku. Każdemu swoje – głosili *Übermensche* ze swastyką i doczekali się szybciej nagrody niż się spodziewali, nie musieliśmy czekać zapowiadanego tysiąca lat. Nie przeciwstawiamy się więc złu, poczekajmy, a wszystko przeminie.

Miał muzykalną żonę, dwie piękne córki – jedna żyje ekonometrią, a druga muzyką. Żona – chociaż muzyk – obdarzona była wielką intuicją ekonomiczną. W okresie dokuczliwej, sporej inflacji z lat osiemdziesiątych XX wieku wślawiła się powiedzeniem: *pieniądz wydany, to pieniądz uratowany*. Żona zmarła kilka lat temu. Dane mi było być w mieszkaniu Państwa Hożerów i wysłuchać Jej kameralnego koncertu; grała na fortepianie dla swej rodziny, profesora Jana Zawadzkiego i dla mnie. Było to w czasach, gdy mieszkał jeszcze w bloku; później wybudował dom z oczkiem wodnym w ogrodzie. Wiem to z opowieści, bo nad tym oczkiem sam nie byłem; oglądając piękne grzybieńce – popularnie zwane liliami wodnymi – profesor Hozer rozmyśla o czasie, jego przemijaniu i wpływie tego czasu na naukę i ekonometrię w szczególności. W Zakopanem propagował napój z francuska zwany *panaché*; było to piwo zmieszane z sokiem jabłkowym. Jeżeli sok i piwo są wysokiej klasy, to *panaché* smakuje wybornie. Dobry humor w połączeniu z żartem – to przepis na codzienne drobne kłopoty. W 2005 roku profesor Andrzej Barczak wymówił się nawalem ważnych prac od mikroekonometrycznej konferencji w Świnoujściu. Jednakowoż podczas konferencji ktoś przypadkowo

spotkał go w tym mieście i przyprowadził na miejsce obrad. – *Wpadłem tu do żony* – usprawiedliwiał się Andrzej. – *Wpadł jak śliwka w kompot* – zamknął sprawę Józef. Na słynnych sympozjach naukowych organizowanych co dwa lata w Świnoujściu je się ryby; tu także powstała *teoria szwagra*. Jest to odmiana głośnej metody znanej pod nazwą *benchmarking* – jeśli musisz się wzorować na kimś to wzoruj się na najlepszym. Szwagier jest przypuszczalnie najlepszy w otoczeniu rodzinnym, więc za sprawą żony jest wzorem do naśladowania. Teorię szwagra zaproponował profesor Ryszard Antoniewicz. Profesor Hozer lubi i ceni fizyków, bo słusznie uważa, że nauka jest uniwersalna i prawa stworzone w jednej dyscyplinie można przez analogię przenieść do innej. Teoria szwagra jest takim właśnie homomorfizmem naukowym w żartobliwym ujęciu. Sam jest także autorem reguły pragmatycznej zwanej *renta pieczątki*. Było to w czasie, gdy urzędnicy państwowi i profesorowie dysponowali jedynie charyzmatem posiadanych uprawnień i wydawali swą akceptację wtedy tylko, gdy sprawa była do załatwienia: *jak się da, to się robi*. Propaguje i popularyzuje naukę w różny sposób – także przez celne wypowiedzi dla mediów. Ekonometria – panna o raczej słabym powodzeniu – musi być teraz silnie promowana. Kiedyś było inaczej, był to kierunek najbardziej oblegany. Wierzono bowiem, że matematyka jest panaceum na wszystkie socjalistyczne przemijające trudności. Rozczarowanie przyszło w swoim czasie, jednakowoż zaakcentować trzeba mocno, że bez matematyki nie ma nauki.

Chociaż jest człowiekiem o silnym charakterze stosuje metody *soft power* w oddziaływaniu na kolegów i współpracowników. Brzydzi się kłamstwem. Ongiś – w okresie stanu wojennego – pędził bez prawa jazdy autostradą amerykańską z niedozwoloną szybkością, pożyczonym od kolegi samochodem. Zatrzymuje go policja i jak zwykle żąda papiery. Kierowca przyznaje się do wszystkiego; swą szczerą spowiedzią wzbudził takie zaufanie, że strażnicy porządku drogowego życzyli mu szczęśliwego pobytu w Stanach i nie było mowy o jakiegokolwiek karze za występki. Było to w czasach, gdy Ameryka darzyła wielką sympatią Polaków z powodu *Solidarności* i naszej walki z sowiecką zależnością, z socjalistycznymi porządkami. Ma szczęśliwą rękę, jest człowiekiem, który wszystko ulepszy, gdy tylko przyłoży swą rękę do sprawy. Nie pokłada jednak wartości w dobrach materialnych. Jest kierownikiem katedry, profesorem zwyczajnym, redaktorem *Przeglądu Statystycznego*, członkiem Komitetu Statystyki i Ekonometrii PAN, dyrektorem Instytutu Diagnoz Ekonomicznych *et cetera, et cetera*. Jest człowiekiem, który mógłby zrobić wielką karierę w polityce lub biznesie, gdyby tylko chciał. *Gdy bogactwo się mnoży, nie lgnijcie do niego sercem* (Ps 62. 11). Na dworcu w Krakowie kieszonkowi złodzieje zrobili wokół niego ongiś sztuczny ścisk i zabrali jego portfel z pokaźną gotówką. Niewątpliwie zdawał sobie sprawę z tego, co się dzieje, więc wykrzyknął: – *Zwróćcie chociaż dokumenty!* Istotnie po kilku dniach otrzymał w przesyłce portfel z pełną zawartością, a tylko – błałostka – bez pieniędzy. Zdarzenie to jest dla niego dowodem prawdziwej etyki złodziejskiej, etyki zawodowej. Cechuje go miłość wszystkich stworzeń – postawa konstruktywna, pomoc. Jest daleki od działań destruktywnych, od szkodenia i psucia. Brzydzi się kłamstwem i rzucaniem błotem na tej zasadzie, że zawsze coś przylgnie. *Calumniare audacter, semper aliquit haeret* – szkaluj zuchwale, zawsze coś zostanie. Chrześcijańskie miłosierdzie jest dla niego codzienną strawą.

6. Teoria kategorii jest szczytowym osiągnięciem matematyki. Nie istnieje zbiór wszystkich zbiorów, nie ma rodziny wszystkich grup, nie ma klasy wszystkich przestrzeni liniowych, ale są kategorie. Jest kategoria mnogości, kategoria grupy abelowej, kategoria przestrzeni liniowej itd. Kategorie wywodzą się od Arystotelesa, który wyróżnia dziesięć ich rodzajów: *substantia, quantitas, qualitas, relatio, actio, passio, ubi, quando, situs, habitus*. Profesor Hozer wymienia te kategorie od czasu do czasu, by przypomnieć młodym pracownikom nauki, że ten układ może być wielce pomocny przy planowaniu wszelkich badań. Schematem tym posługują się dziennikarze w opisie wszystkich wydarzeń; komunikat dziennikarski, podobnie jak artykuł naukowy, winien wskazywać czytelnikowi jasno co się zdarzyło, gdzie i kiedy, z jakiego powodu. Jest zwolennikiem teleologii – nauki o celowości w przyrodzie. Teleologia pozwala lepiej zrozumieć przyrodę i jest jednocześnie nadrzędnym prawem uzasadniającym wszystkie prognozy. Znajomość celu pozwala zrozumieć postępowanie i zobaczyć najkrótszą drogę. Jedynym celem przyrody, gospodarki i społeczeństwa jest równowaga. Zakłócenie równowagi destabilizuje system. Mrowisko jest stabilne, bo jego struktura jest stała. Ludzie lubią odmianę; ta nasza ciekawość nowości bywa przyczyną nieszczęść. *A ja wiem, że w życiu jest równowaga. Boję się, że zapłacę za nadmiar. Staram się zachować proporcje – znakomitą whisky zawsze zagryzam tanią parówką* [8]. Są to słowa dwudziestopięcioletniej aktorki, której wszystko się udaje pewnie dlatego, że jest zdolna i ciężko pracuje. Natura jest celowa; uprzykrzone komary są niezbędne dla równowagi globalnej, bo są pokarmem jaszczurek, żab i wielu innych płazów. Krety z kolei są nie lubianymi gośćmi w ogrodach, ale ich działanie jest bardzo celowe. Spulchnienie ziemi sprzyja wzrostowi roślin i pomaga w rozmnażaniu dżdżownic, które są pokarmem również kreta. To krecia praca była przypuszczalnie wzorem dla pierwszych rolników; orka jest ucywilizowanym sypaniem kopców przez kreta.

Liczne referaty profesora Hożera, które miałem przyjemność wysłuchać charakteryzują się zawsze dwoma zasadniczymi cechami. Są krótko i jasno przedstawione – podobnie jak wystąpienia profesora Michała Kolupy, czyli należą do kategorii *non multa, sed multum* – oraz pełne treści. Czas jest według niego jedynym czynnikiem sprawczym, uniwersalną przyczyną. Struktura zmienia się z czasem, a trwanie jest tożsame z rozwojem. Nie ma czasu bez zmiany; czas jest następstwem zdarzeń. Stabilność nie jest stagnacją. Jest oznaką stałej struktury w dążeniu do swego przeznaczenia – celu. Zdarzyło się, że na jednej ze wspomnianych konferencji mikroekonometrycznych mówiłem o stabilnej strukturze szczecińskiego zespołu badawczego. Moje wystąpienie zrozumiała opacznie pani profesor Elżbieta Żółtowska z Łodzi i pouczyła mnie, że zespół ten daleki jest od stagnacji i dynamicznie się rozwija. Stabilność nie ma nic wspólnego ze stagnacją, ale oznacza, że małe zakłócenia nie wytrącają z obranej drogi. Jest to równowaga stabilna, a stagnacja jest także równowagą, ale zwykle niestabilną. Struktura jest zależna od czasu. Bifurkacja trajektorii oznacza zmianę struktury. Struktura jest naturalnie wyznaczona przez relację pomiędzy wybranymi wielkościami, struktura jest pochodną prawa przyrody będącego wspomnianą relacją. *Relacje, czyli prawa nauki, są gładkimi różniczkami w przestrzeni stanów*. Oczywiście, zawsze pierwsza pojawia się przyczyna a później skutek. Ponieważ czas definiuje się przez następstwo zdarzeń, więc można powiedzieć, że czas jako powszechna przyczyna jest rezultatem skutku; skutek wyprzedza przyczynę. *Sic!* Przy okazji przypomnę dowcipną i mądrą anegdotę,

którą posłużył się profesor Teodor Kulawczuk w dyskusji nad starym zagadnieniem związku przyczynowego. Jedyny znany przypadek, gdy skutek poprzedza przyczynę, to udział lekarza w kondukcje pogrzebowym swego pacjenta. Czas wszystko zmienia; nie wchodzi się dwa razy do tej samej rzeki uczyli dawno temu Grecy.

7. Z prawem pięciu gospodarstw domowych na jedno przedsiębiorstwo związana jest liczba złota χ . Powszechnym prawem nauki jest zasada równowagi. Z tej zasady wynika ruch wirowy; ruch wirowy jest najczęściej spotykaną odmianą rozwoju. Ruch wirowy zawiera w sobie liczbę złotą, bo w tym ruchu jest spirala logarytmiczna. Pokażemy, że liczba złota jest tak samo wielkością ekstremalną, czyli prawem nauki, jak liczby π oraz e . Te prawa przyrody to najpiękniejszy przejaw ogólnej zasady równowagi. Równowaga wymusza porządek, harmonię i symetrię. Nawet system obrachunkowy jest praktyczny wtedy, gdy jest oparty na zasadzie równowagi – dużo partycji i mało dobrze rozróżnialnych nominałów. Podobnie jest z ergonomią i wszystkimi konstrukcjami technicznymi usprawniającymi pracę i życie. Zwykły układ nakrętki i przeciwnakrętki stabilizuje połączenie, albowiem pracuje w systemie wzajemnej równowagi. Liczba złota

$$\chi = \frac{\sqrt{5} - 1}{2}$$

dzieli odcinek $[0, 1]$ w stosunku harmonicznym, czyli spełnia równanie złotego cięcia

$$\frac{x}{1} = \frac{1-x}{x}.$$

Jest to liczba algebraiczna – niewymierna – równa 0.6180..., zwana boską proporcją, która rządzi pięknem. Spotykamy ją w architekturze i w przyrodzie. Doświadczenie jest podstawą kanonu piękna odwołującego się do złotego podziału. Złota liczba też została wyznaczona z próby. Wielomian złotego podziału $s(x) = x^2 + x - 1$ jest oczywiście nierozkładalny w ciele liczb wymiernych Q , lecz jest rozkładalny w ciele $Q(\sqrt{5})$. Złota liczba χ jest granicą ułamka łańcuchowego $1/(1 + 1/(1 + 1/(1 + \dots)))$. Łatwo sprawdzić, wynika to z ogólnego kryterium, że ułamek ten jest zbieżny. Jeżeli jest to rozwinięcie w ułamek łańcuchowy liczby x , wtedy $x = 1/(1 + x)$. Oznacza to, że ten najprostszy i najpiękniejszy ułamek łańcuchowy reprezentuje liczbę złotą. Każdy wielomian $p(x) = ax^2 + bx + c$ o współczynnikach całkowitych, $a, b, c \in Z$, będziemy utożsamiać z uporządkowaną trójką (a, b, c) . Zbiór takich wielomianów jest modułem Z^3 nad pierścieniem liczb całkowitych Z . W zbiorze Z^3 wyróżniamy podzbiór F wielomianów stopnia drugiego – oznacza to, że współczynnik a jest różny od zera, o pierwiastkach rzeczywistych, czyli $\Delta(a, b, c) = b^2 - 4ac \geq 0$, nierozkładalnych nad ciałem Q , więc wielkość $\Delta(a, b, c)$ nie jest kwadratem liczby całkowitej. Tak więc $(a, b, c) \in F$ wtedy, i tylko wtedy, gdy po pierwsze, $a \neq 0$, po drugie, wyróżnik nie jest ujemny, $b^2 - 4ac \geq 0$, i po trzecie, $\sqrt{b^2 - 4ac}$ nie jest liczbą wymierną. Funkcja $\Delta : Z^3 \rightarrow Z$ na zbiorze F przyjmuje wartości nieujemne. Istnieje więc taki wielomian $(a^*, b^*, c^*) \in F$, że $\Delta(a^*, b^*, c^*) = \min \Delta(F)$, gdzie $\Delta(F) = \{\Delta(a, b, c) : (a, b, c) \in F\}$. Wielomian spełniający powyższy warunek nazywamy wielomianem minimalnym.

Lemat. Wielomian złotego podziału $(1, 1, -1)$ jest wielomianem minimalnym oraz

$$\Delta(1, 1, -1) = \min \Delta(F) = 5.$$

Zbiór wielomianów minimalnych jest oczywiście nieskończony. Można go utożsamić ze zbiorem rozwiązań równania diofantycznego $b^2 - 4ac = 5$. Każdy wielomian stopnia drugiego (a, b, c) rozkładalny nad ciałem $\mathbb{Q}(\sqrt{5})$ i nierozkładalny nad ciałem $\mathbb{Q}$ ma wyróżnik $\Delta(a, b, c)$ postaci $2^{2k}(2n + 1)^2 5$, gdzie $k, n \in \mathbb{N}$.

Dowód. Warunek

$$\Delta(a, b, c) = 2^{2k}(2n + 1)^2 5$$

jest konieczny i wystarczający na to, aby wielomian $p(x) = ax^2 + bx + c$ był nierozkładalny w ciele $\mathbb{Q}$ i rozkładalny w ciele $\mathbb{Q}(\sqrt{5})$. Widać stąd, że minimalna wartość funkcji Δ na zbiorze F jest wtedy 5, gdy $k = n = 0$. Jeżeli $\Delta(a, b, c) = d^2$, to oczywiście wielomian (a, b, c) nie należy do zbioru F . Oznacza to, że liczby naturalne 0, 1 i 4 nie są w zbiorze wartości $\Delta(F)$ funkcji Δ na zbiorze F . Przypuśćmy, że dla pewnego wielomianu $(a, b, c) \in F$ jest $\Delta(a, b, c) = 2$, czyli że $b^2 - 4ac = 2$. Wynika stąd, że b jest liczbą parzystą, $b = 2r$, $r \in \mathbb{Z}$, a więc $4r^2 - 4ac = 2$. Równość ta jest jednak niemożliwa, bo lewa strona jest podzielna przez 4, a prawa nie jest. Podobnie sprawdza się, że 3 nie jest wartością funkcji Δ na F . Jeżeli $b^2 - 4ac = 3$ dla pewnego $(a, b, c) \in F$, to b nie może być liczbą parzystą, czyli jest postaci $2r + 1$. Jest więc $4r^2 + 4r - 4ac = 2$, ale ta równość też jest niemożliwa. Podane argumenty dowodzą, że funkcja Δ osiąga minimalną wartość 5 nie tylko na zbiorze wielomianów stopnia 2 nierozkładalnych nad $\mathbb{Q}$ i rozkładalnych nad $\mathbb{Q}(\sqrt{5})$, ale także na zbiorze F dowolnych wielomianów stopnia 2 nierozkładalnych nad $\mathbb{Q}$ i rozkładalnych nad $\mathbb{Q}(\sqrt{d})$, gdzie liczba d , zależna od wielomianu, jest liczbą naturalną niebędącą pełnym kwadratem. Dowód został skończony.

Liczba złota jest uniwersalnym prawem przyrody. Dodatkowym argumentem jest ciąg Fibonacciego: $a_0 = 1$, $a_1 = 1$, $a_{n+2} = a_{n+1} + a_n$, $n \in \mathbb{N}$. Ciąg stosunków $\left(\frac{a_{n+1}}{a_n}\right)$ jest zbieżny, jego granicą jest liczba $\frac{1+\sqrt{5}}{2}$ – odwrotność liczby złotej. Oznacza to, że ciąg Fibonacciego asymptotycznie zachowuje się jak ciąg geometryczny o ilorazie $q = \frac{1}{\chi}$. Króliki rozmnażają się więc zgodnie z prawem Malthusa. Zapis liczby złotej w postaci ułamka łańcuchowego pokazuje, jak prosty, piękny i naturalny jest to obiekt; wielkości $\frac{a_{n+1}}{a_n}$ są reduktami tego ułamka.

8. Wspomniana wyżej piękna książka z teorii liczb – powstała w katedrze profesora Hozera – była oceniana przez wybitnego specjalistę z tej dziedziny – profesora Andrzeja Schinzla. Po ukazaniu się dzieła w druku wielki recenzent okazał swą mało-
 stkowość, bo wdał się w śmieszny polemikę na temat swej własnej opinii. Wspominam

o tym, by zaznaczyć, że życie codzienne ludzi nauki obfituje w kłopotliwe drobiazgi, a jednocześnie pragnę wskazać na fundamentalne znaczenie teorii liczb w ekonomii i całej nauce. Nauka o podzielności liczb algebraicznych jest dziedziną *par excellence* ekonomiczną. Podzielność – to relacje porządku i preferencji. Pitagoras głosił, że liczby rządzą światem. Tezę tę adoptował profesor Hozer. Sprzeciwił się jednak temu pewien rygorystyczny kapłan, który słusznie uważa, że to tylko Bóg rządzi światem. Mądrość Boża jest podstawą wszelkiego stworzenia; prawa boskie są prawami natury. Tych właśnie praw szuka nauka; ich istotą są proporcje i harmonia. Proporcje rodzą liczby; liczba złota jest właśnie boską harmonią. Zdarzenia cudowne z definicji przeczą prawom nauki, ale są i takie zdarzenia cudowne, które nauka potrafi uzasadnić. Przykładem niech będzie opowieść biblijna o przepiórkach. Bóg wzbudził wiatr południowy, który przyniósł obfitość mięsa w postaci przepiórek na zgłodniały obóz wędrowców do ziemi obiecanej. Zdarzenie to można interpretować jako całkowicie naturalne. I tak właśnie niektórzy mędrcy współcześni czynią. Jeśli mówimy o wyjściu z Egiptu, to koniecznie wspomnieć trzeba o Józefie. Był piękny, mądry i co najważniejsze – był wybrańcem Boga. Widzę podobieństwo Józefa Hozera do swego biblijnego imiennika. *The instruction he gave was true and no word of injustice fell from his lips* (Mal. 2.6).

Profesor Hozer ceni i lubi piękno, szczególnie delectuje się malarstwem. U niego w domu widziałem dwa obrazy profesora Andrzeja Alexiewicza z Poznania. Obrazy tego autorstwa spotkałem także u nieżyjącego już profesora Eugeniusza Szczepankiewicza. Profesor Alexiewicz słynął z prawości, uczoności i dobrych manier; był przyjacielem lwowskim profesora Rudolfa Hohenberga – trzeciego w kolejności kierownika Katedry Matematyki Akademii Ekonomicznej we Wrocławiu, a także przyjacielem ucznia profesora Hohenberga – profesora Szczepankiewicza. Alexiewicz jest autorem encyklopedycznego działu z analizy funkcjonalnej i słynnego powiedzenia: *Lepiej być panem siebie, niż członkiem PANu*. Ta wypowiedź jest, oczywiście, żartobliwą oceną wiodącej instytucji naukowej w Polsce. W akademii mamy dużo członków, a mało prawdziwych uczonych. Wszelkie dywagacje naukowe zawsze z konieczności zahaczają o akademię. Profesor Józef Hozer jest naturalnie panem siebie; jest możliwe, że stanie się członkiem PANu. Wskazuje na to jego dorobek naukowy, działalność organizacyjna oraz nieprzeciętne zalety charakteru. Życzę mu tego gorąco, wszak byłby to wielki sukces całego środowiska.

To, co się zdarzyło jest właśnie konieczne, nie można było tego uniknąć; to, co jest możliwe staje się konieczne tylko wtedy, gdy się zrealizuje. Nie planowałem wyjazdu do Świnoujścia w 2009 roku z wielu powodów. Widocznie jest koniecznością abym tam był, bo nie mogłem odmówić osobistej prośbie organizatora. Mogę jedynie zawołać za świętym Pawłem: *I am an unprofitable servant; I have done that which was my duty to do*. Uczyniłem co było konieczne. Wierzę głęboko w przeznaczenie i mądrość Bożą, która kieruje naszą ręką.

Warszawa, 16 lipca 2009.

LITERATURA

- [1] Cieśliński P., [2009], *Naukowe dowody na fałszerstwo w Iranie*, „Gazeta Wyborcza” z 27 czerwca 2009, s. 2.
- [2] Hozer J., [1993], *Mikroekonometria: analizy, diagnozy, prognozy*, Państwowe Wydawnictwo Ekonomiczne, Warszawa.
- [3] Hozer J., [1998], *Nieruchomości, przedsiębiorstwa, wyceny, analizy*, Pomoc i Rozwój, Szczecin.
- [4] Hozer J., [1998a], *Czas i przestrzeń w modelowaniu ekonometrycznym, czyli tempus, locus, homo, casus et fortuna regit actum*, Wydawnictwo Akademii Ekonomicznej w Krakowie.
- [5] Hozer J., Doszyń M., [2004], *Ekonometria skłonności*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- [6] Hozer J., Markowicz I., [2002], *Małe firmy: analizy i diagnozy*, Wydawnictwo Uniwersytetu Szczecińskiego.
- [7] Smoluk A. (red.), [2000], *Elementy metrologii ekonomicznej*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu.
- [8] Żmuda M., [2009], *Myślę, mówię, mam!*, Twój Styl 8, s. 40-44.

PODZIĘKOWANIE RECENZENTOM

Redakcja pragnie serdecznie podziękować recenzentom za podjęty trud opiniowania opracowań naukowych nadesłanych w 2009 r. do „Przeglądu Statystycznego”. W wielu przypadkach opinie recenzentów pozwoliły na znaczne udoskonalenie opracowań przed ich publikacją. Dziękujemy za wnikliwość, staranność i terminowość recenzji.

W kończącym się 2009 r. podziękowanie to składamy następującym osobom:

Dr hab. Tadeusz Bołt,
Prof. dr hab. Wojciech Charemza,
Prof. dr hab. Piotr Chrzan,
Prof. dr hab. Czesław Domański,
Dr hab. Eugeniusz Gatnar,
Dr hab. Marek Gruszczyński,
Prof. dr hab. Krzysztof Jajuga,
Prof. dr hab. Stanisław Maciej Kot,
Dr Grzegorz Krzykowski,
Prof. dr hab. Mirosław Krzyśko,
Prof. dr hab. Władysław Milo,
Dr hab. Paweł Miłobędzki,
Dr Krzysztof Najman,
Prof. dr hab. Jacek Osiewalski,
Prof. dr hab. Magdalena Osińska,
Prof. dr hab. Emil Panek,
Dr hab. Mariola Piłatowska,
Dr hab. Mateusz Pipień,
Prof. dr hab. Józef Stawicki,
Dr hab. Krystyna Strzała,
Prof. dr hab. Bogdan Suchecki,
Prof. dr hab. Marek Walesiak.

SPRAWOZDANIA

BARBARA BATÓG, IWONA MARKOWICZ

XIV OGÓLNOPOLSKA KONFERENCJA „MIKROEKONOMETRIA W TEORII I PRAKTYCE”

W dniach 3-5 września 2009 roku odbyła się XIV Ogólnopolska Konferencja „Mikroekonometria w teorii i praktyce”. Konferencja została zorganizowana przez Katedrę Ekonometrii i Statystyki Uniwersytetu Szczecińskiego oraz Instytut Analiz, Diagnoz i Prognoz Gospodarczych w Szczecinie. Obrady odbyły się na promie „Pomerania” w drodze ze Świnoujścia do Kopenhagi. Konferencja była również okazją do uczczenia 40-lecia pracy naukowej profesora Józefa Hozera, kierownika Katedry Ekonometrii i Statystyki.

W konferencji wzięło udział około 100 osób z 21 ośrodków naukowych z całej Polski oraz z praktyki. Zostało wygłoszonych 17 referatów i zaprezentowano około 60 posterów. Po owocnych obradach i dyskusjach naukowych uczestnicy konferencji wzięli udział w wycieczce po Kopenhadze.

Poniżej prezentujemy streszczenia wygłoszonych referatów (w kolejności alfabetycznej). Z treścią wygłoszonych referatów i zaprezentowanych posterów będzie można zapoznać się po wydaniu przez Uniwersytet Szczeciński recenzowanej monografii.

Prof. dr hab. Maria Balcerowicz-Szkutnik, prof. dr hab. Elżbieta Sojka (Akademia Ekonomiczna w Katowicach), *Analiza aktywności ekonomicznej osób w wieku 50+ w wybranych krajach UE*

Autorki przedstawiły ogólne tendencje zmian w strukturze ludności według wieku oraz w aktywności zawodowej pokolenia 50+ w krajach nadbałtyckich Unii Europejskiej. Cztery z tych krajów, tj. Dania, Finlandia, Niemcy, Szwecja, to kraje „starej” Unii natomiast Estonia, Litwa, Łotwa i Polska należą do grupy nowych członków UE. Badaniem objęto lata 1996-2008 oraz przedstawiono prognozy liczby osób w wieku 50+. Dokonano również analizy stanu obecnego i prognozy zmian parametrów charakteryzujących rynek pracy, czyli współczynników aktywności zawodowej i wskaźnika zatrudnienia oraz stopy bezrobocia ogółem i bezrobocia długoterminowego.

Prof. dr hab. Stanisława Bartosiewicz (Uniwersytet Ekonomiczny we Wrocławiu), *Magiczne słowo „produkt”*

Autorka przedstawiła kilka refleksji dotyczących terminologii naukowej, a w szczególności używania teraz terminu (zdefiniowanego dawno) w znacznie rozszerzonym zakresie. Termin ten to właśnie wymienione w tytule magiczne słowo produkt. Rozważone zostały skutki wynikające z rozszerzenia pojęcia *produkt* na pojęcie *usługa* w procedurze konstrukcji funkcji produkcji (zmiany w mierzeniu produkcji oraz czynników warunkujących jej zmienność) oraz w trudnościach wyróżnienia roli pracy żywej świadczonej na rzecz usług.

Prof. dr hab. Andrzej Bąk (Uniwersytet Ekonomiczny we Wrocławiu), *Zastosowanie mikroekonometrycznych modeli zmiennych dychotomicznych w badaniach preferencji i ich estymacja w programie R*

Program R zawiera pakiety i funkcje, które można wykorzystać w celu analizy mikrodanych o preferencjach i wizualizacji wyników. Autor przedstawił przykład analizy mikrodanych ankietowych za pomocą modeli zmiennych dychotomicznych szacowanych w programie R. Mikrodane dotyczyły preferencji w zakresie odżywiania się studentów oraz preferencji osób urządzających ogrody.

Prof. dr hab. Ewa Drabik (Szkoła Główna Gospodarstwa Wiejskiego), *O pewnej modyfikacji teorii fal Elliotta*

Teoria fal Elliotta to jedno z ciekawszych narzędzi diagnostycznych odnoszącym się do prognozowania przyszłych ruchów cen na giełdzie. Celem pracy było sprawdzenie, w jaki sposób modyfikacja ciągu Fibonacciego wprowadzona przez Ulama, zawierająca element losowości, pozwoli na bardziej adekwatne do rzeczywistości analizowanie fal giełdowych. Wprawdzie ze względu na nieskończoną liczbę realizacji tego ciągu prognozowanie jest raczej trudne, ale być może w przyszłości uda się również i ten problem rozwiązać.

Prof. dr hab. Jan Gajda (Uniwersytet Łódzki), *Ekonometryczna wycena wartości mieszkania*

Autor zbudował model ekonometryczny służący do szacowania wartości nieruchomości. Rozpatrywano duży zestaw potencjalnych zmiennych objaśniających zarówno ilościowych, jak i jakościowych opisujących cechy fizyczne mieszkań z bazy danych liczącej kilka tysięcy obiektów. Oprócz tego typu zmiennych, do potencjalnych zmiennych objaśniających, zaliczono również kilka zmiennych z otoczenia makroekonomicznego.

Prof. dr hab. Stefan Grzesiak, mgr inż. Robert Józwiak (Uniwersytet Szczeciński), *Wybór miejsca i formy opodatkowania firmy jako problem decyzyjny*

Autorzy skupili uwagę na problematyce dotyczącej małych i średnich firm, które mają z reguły najwięcej kłopotów z otoczeniem ekonomiczno-prawnym, a jednocześnie z braku dostatecznego potencjału ekonomicznego nie posiadają zorganizowanej własnej bazy doradczej. Zasadniczym tematem prezentowanej pracy był wybór przez przedsiębiorców zamierzających prowadzić lub już prowadzących działalność gospodarczą w Polsce i w Niemczech miejsca i sposobu opodatkowania. Ze względu na historyczną odmienną systemów podatkowych w obu krajach, przedsiębiorca powinien świadomie dokonywać wyboru kraju, w którym będzie zarejestrowany i w którym powinien opłacać podatek. Dla ogólności rozważań założono, że podatek jest płacony w obu krajach, a efektem przeprowadzonej dyskusji i w następstwie procedury optymalizacyjnej powinien być taki dobór dochodów w obu krajach, aby suma płaconych podatków była minimalna.

Prof. dr hab. Józef Hozier (Uniwersytet Szczeciński), *Sklonności a teoria uczuć i działań Vilfredo Pareto*

Autor odniósł się do wprowadzonego przez Pareto pojęcia rezyduów, zwanych też osadami psychicznymi. Podział rezyduów na klasy według Pareto można modyfikować. Należałoby skłonności podzielić według zasadniczych celów działań ludzkich, które ujmuje geometryczna interpretacja osobowości ludzkiej: cele materialne, intelektualno-duchowe i seksualne. Na tym tle Autor prowadził rozważania o skłonnościach, które odgrywają znaczącą rolę w przestrzeni ekonomicznej.

Prof. dr hab. Michał Kolupa (Politechnika Radomska), *O wyznaczaniu pewnych charakterystyk jednorównaniowego liniowego modelu ekonometrycznego*

W pracy omówiono wyznaczanie zarówno wektora wartości teoretycznych, wektora reszt oraz wartości prognozowanej zmiennej endogenicznej jednorównaniowego modelu ekonometrycznego szacowanego klasyczną metodą najmniejszych kwadratów. Przedstawiono zarówno klasyczne postępowanie, jak i postępowanie wykorzystujące macierze brzegowe. W praktyce ekonometrycznej bardzo często zachodzi konieczność wyznaczenia niecałego wektora wartości teoretycznych, jak również niecałego wektora reszt, lecz jedynie wybranej jego składowej. W tym przypadku powinno być zastosowane twierdzenie Kroneckera-Capelliego.

Prof. dr hab. Anna Malina (Uniwersytet Ekonomiczny w Krakowie), *Ocena realizacji założeń i celów Strategii Lizbońskiej w obszarze rynku pracy w krajach Unii Europejskiej*

Autorka przedstawiła założenia i cele Strategii Lizbońskiej w dziedzinie zatrudnienia. Skoncentrowała uwagę na dyrektywach dotyczących wskaźników zatrudnienia w UE (zatrudnienie ogółem, kobiet oraz starszych pracowników). Przedstawiła także sytuację w zakresie bezrobocia w Polsce na tle krajów Unii Europejskiej oraz zwróciła uwagę na dyrektywę dotyczącą czasu pracy. Z przeprowadzonych badań wynika, że sytuacja Polski na tle krajów UE pod względem zatrudnienia jest dość niekorzystna. Dotyczy to zarówno zatrudnienia mężczyzn, kobiet, jak i starszych pracowników.

Prof. dr hab. Paweł Miłobędzki (Uniwersytet Gdański), *Czy ceny akcji na Gieldzie Papierów Wartościowych w Warszawie S.A. są prognozowalne? Uwagi w świetle wyników testów hipotezy błędzenia przypadkowego*

Autor testował hipotezę błędzenia przypadkowego: zmiany cen (instrumentów finansowych) kształtują się w taki sposób, że są nieprognozowalne. Otrzymane wnioski to: brak podstaw do odrzucenia hipotezy błędzenia przypadkowego dla wszystkich głównych indeksów GPW w Warszawie S.A. oraz jej modyfikacji (niesymetryczny powrót do średniej logarytmicznych stóp zwrotu), stopy zwrotu z większości portfeli naśladowujących główne indeksy giełdowe GPW w Warszawie S.A. – poza stopami zwrotu z portfeli naśladowujących indeksy WIG20 i WIG-Telekomunikacja – są prognozowalne, pomimo prognozowalności stóp zwrotu brak efektywności rynku (rynków cząstkowych) wątpliwy – średnie stopy zwrotu bliskie zeru.

Prof. dr hab. Jerzy Czesław Ossowski (Politechnika Gdańska), *Mikro i makroekonomiczne podstawy zapotrzebowania na pracę w teorii i rzeczywistości gospodarki polskiej*

Autor przedstawił najpierw teoretyczne modele zapotrzebowania na pracę, a następnie zaprezentował oszacowanie dynamicznego kwartalnego modelu zapotrzebowania na pracę w dwóch wariantach. Dokonał również symulacji dynamiki produktu krajowego i dynamiki zapotrzebowania na pracę.

Prof. dr hab. Krzysztof Piasecki (Uniwersytet Ekonomiczny w Poznaniu), *Zbiory intuicyjne w analizie rynku finansowego – możliwości i perspektywy*

Jednym z problemów, z jakimi spotykamy się przy zarządzaniu rynkiem finansowym jest pewna nieprecyzyjność pojęć stosowanych przy analizowaniu jego właściwości. Do opisu tej nieprecyzyjności stosuje się najczęściej teorię podzbiorów rozmytych. Pojęcie zbioru intuicyjnego jest istotnym rozszerzeniem pojęcia zbioru rozmytego i dlatego uzyskiwane tą drogą wyniki mogą wzbogacić naszą dotychczasową wiedzę o rynkach finansowych. W pracy Autor zaproponował metodę modelowania stanów rynku finansowego za pomocą zbiorów intuicyjnych.

Prof. dr hab. Józef Pocięcha (Uniwersytet Ekonomiczny w Krakowie), *Probabilistyczne podejście w rewizji finansowej*

Rewizja sprawozdań finansowych opiera się na uregulowaniach prawnych krajowych, wynikających z ustawy o rachunkowości, krajowych standardach rachunkowości i standardach rewizji finansowej lub uregulowaniach międzynarodowych, wynikających z dyrektyw UE, MSSF czy MSRF. Biegli posługują się także normami wykonywania zawodu i szeregiem procedur rewizyjnych, które w większości mają charakter probabilistyczny. Autor postuluje zwrócenie szczególnej uwagi na dobór reprezentatywnej próby z populacji operacji księgowych w badanym okresie (najczęściej rocznym), ocenę konsekwencji przyjmowania założenia o normalności rozkładu błędów w populacji, uwzględnienie szerszej gamy rozkładów błędów w populacji operacji księgowych oraz poszukiwanie testów najmocniejszych dla badania skuteczności działania systemów kontroli wewnętrznej.

Prof. dr hab. Mirosław Szreder (Uniwersytet Gdański), *O znaczeniu informacji spoza próby w statystyce i ekonometrii*

Autor wskazuje na konieczność efektywnego korzystania w badaniach statystycznych, w szczególności we wnioskowaniu statystycznym, z innych poza próbą badawczą źródeł informacji (w tym informacji *a priori*). Zatem należy włączyć Bayesowskie wnioskowanie do programów nauczania statystyki matematycznej, gdyż istnieją coraz bogatsze zbiory informacji o badanych populacjach, zgromadzonych doświadczeniach i wiedzy badaczy, które mogą poprawić jakość wnioskowania w stosunku do podejścia opartego wyłącznie na próbie losowej, a z drugiej strony model Bayesowskiego wnioskowania, wolny od problemów numerycznych, stanowi bardziej adekwatny od klasycznego podejścia opis przetwarzania przez człowieka informacji i podejmowania decyzji.

Prof. dr hab. Jan Jacek Sztudzynger, dr Paweł Baranowski (Uniwersytet Łódzki), *Rodzinny kapitał społeczny a wzrost gospodarczy – analiza dla Polski i 15 krajów Unii Europejskiej*

Ważnym składnikiem kapitału społecznego jest kapitał więzi rodzinnych – tzw. rodzinny kapitał społeczny. Autorzy postawili hipotezę: rodzinny kapitał społeczny – więzi rodzinne wpływają pozytywnie na wzrost gospodarczy. Im więcej rozwodów, w stosunku do nowo zawieranych małżeństw, tym wolniejszy

wzrost gospodarczy. Weryfikacji tej hipotezy dokonano w oparciu o model ekonometryczny dynamiki PKB, która została uzależniona od współczynnika rozwodów (relacji rozwodów do małżeństw), stopy inwestycji oraz stopy inflacji.

Prof. dr hab. Marek Walesiak, dr Andrzej Dudek (Uniwersytet Ekonomiczny we Wrocławiu), *Odległość GDM dla danych porządkowych a klasyfikacja spektralna*

Autorzy zaproponowali modyfikację metody klasyfikacji spektralnej umożliwiającą jej zastosowanie w klasyfikacji danych porządkowych. W tym celu w typowej procedurze klasyfikacji spektralnej dla danych metrycznych w konstrukcji estymatora jądrowego zastosowana została odległość GDM dla danych porządkowych (odległość GDM2).

Prof. dr hab. Józef Zając (Katolicki Uniwersytet Lubelski Jana Pawła II), *Zagadnienia ekstremalne w teorii regresji*

Autor sformułował problem regresji przy ściśle określonych założeniach. Zaprezentował również dwa sposoby efektywnego wyznaczania funkcji regresji. Rozważania zostały poparte przykładami.

SPIS TREŚCI

Errata do nr 2/2009, tom 56	3
Anna Pajor, Artur Prędkie – Estymacja miernika efektywności technicznej w ramach metody DEA	5
Sabina Denkowska, Kamil Fijorek, Marcin Salamaga, Andrzej Sokołowski – Empiryczna ocena mocy testów dla wielu wariacji	26
Helena Gaspar-Wieloch – Dyskretnie zagadnienie lokalizacyjne dla sieci jednorodnych obiektów	40
Anna Zamajska – Zastosowanie metody DEA w klasyfikacji funduszy inwestycyjnych	51
Andrzej Stryjek – Zastosowanie miar zależności zmiennych losowych oraz kopuli Clayтона i Gumbel-Hougaarda do szacowania wartości zagrożonej	67
Michał Kolupa, Joanna Plebaniak – O eliminacji błędów w prognozach wykonywanych na podstawie modelu Leontiewa w przypadku agregacji niedoskonałej w sensie Hatanaki	81
Antoni Smoluk – 40 lat działalności naukowej profesora Józefa Hozera, czyli o metrologii ekonomicznej	89
Podziękowanie Recenzentom	103

SPRAWOZDANIA

Barbara Batóg, Iwona Markowicz – XIV Ogólnopolska Konferencja „Mikroekonometria w teorii i praktyce”	104
--	-----

CONTENTS

Anna Pajor, Artur Prędkie – Estimation of the Technical Efficiency Measure Within the DEA Method	5
Sabina Denkowska, Kamil Fijorek, Marcin Salamaga, Andrzej Sokołowski – A Comparative Study of Tests Power for Homogeneity of Variances	26
Helena Gaspar-Wieloch – The Discrete Location Problem for a Chain of Homogeneous Facilities	40
Anna Zamajska – Efficiency of Mutual Fund and a Data Envelopment Analysis	51
Andrzej Stryjek – Application of Random Variables Dependence Measures and Clayton and Gumbel-Hougaard Copulas for Estimating Value at Risk	67
Michał Kolupa, Joanna Plebaniak – On Eliminating Errors of Forecasts Obtained from the Leontief Model for the Case of Imperfect Hatanaki's Aggregation	81
Antoni Smoluk – The Forty Years Professor's Józef Hozer Scientific Activity on Economic Metrology	89
Acknowledgements	103

REPORTS

Barbara Batóg, Iwona Markowicz – The XIV Nationwide Conference „Microeconometrics in Theory and Practice”	104
---	-----