

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 64

2

2017

WARSZAWA 2017

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA PROGRAMOWA

Andrzej S. Barczak, Czesław Domański, Marek Gruszczyński, Krzysztof Jajuga (Przewodniczący),
Tadeusz Kufel, Igor G. Mantsurov, Jacek Osiewalski, Sven Schreiber, Mirosław Szreder, D. Stephen
G. Pollock, Jaroslav Ramík, Peter Summers, Matti Virén, Aleksander Welfe, Janusz Wywiał

KOMITET REDAKCYJNY

Magdalena Osińska (Redaktor Naczelny)
Marek Walesiak (Zastępca Redaktora Naczelnego, Redaktor Tematyczny)
Michał Majsterek (Redaktor Tematyczny)
Maciej Nowak (Redaktor Tematyczny)
Anna Pajor (Redaktor Statystyczny)
Piotr Fiszedler (Sekretarz Naukowy)

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

Strona WWW „Przegląd Statystyczny”:
<http://www.przegladstatystyczny.pan.pl>

Nakład 140 egz.



Przygotowanie do druku:
Dom Wydawniczy ELIPSA
ul. Inflancka 15/198, 00-189 Warszawa
tel./fax 22 635 03 01, 22 635 17 85
e-mail: elipsa@elipsa.pl, www.elipsa.pl

Druk:
Centrum Poligrafii Sp. z o.o.
ul. Łopuszańska 53
02-232 Warszawa

SPIS TREŚCI

Emil P a n e k – Gospodarka Gale’a z wieloma magistralami. „Silny” i „bardzo silny” efekt magistrali.	137
Andrzej S o k o ł o w s k i, Małgorzata M a r k o w s k a – Iteracyjna metoda liniowego porządkowania obiektów wielocechowych.	153
Karol K u k u ł a, Lidia L u t y – Jeszcze o procedurze wyboru metody porządkowania liniowego	163
Bartłomiej L a c h – Metody łączenia i selekcji klasyfikatorów w prognozowaniu upadłości przedsiębiorstw.	177
Grzegorz T a r c z y ń s k i – Model dla zadania kompletacji zleceń łączonych uwzględniający problem blokowania się magazynierów	193
Jakub W i t k o w s k i – Zastosowanie metod badań operacyjnych w ocenie efektywności mechanizmów selekcji w organizacjach charytatywnych	213

CONTENTS

Emil P a n e k – Gale’s Economy with Multiple Turnpikes. “Strong” and “Very Strong” Turnpike Effect	137
Andrzej S o k o ł o w s k i, Małgorzata M a r k o w s k a – Iterative Method for Linear Ordering of Multidimensional Objects.	153
Karol K u k u ł a, Lidia L u t y – Once More about the Selection Procedure for the Linear Ordering Method.	163
Bartłomiej L a c h – Classifier Combination and Selection Methods at the Corporate Bankruptcy Prediction.	177
Grzegorz T a r c z y ń s k i – Model for Order-Batching with Congestion Consideration.	193
Jakub W i t k o w s k i – Assessing Selection Mechanisms in Charity Organizations Using Operational Research Methods	213

EMIL PANEK¹GOSPODARKA GALE’A Z WIELOMA MAGISTRALAMI.
„SILNY” I „BARDZO SILNY” EFEKT MAGISTRALI

1. WSTĘP

W artykule Panek (2016b) przedstawiliśmy model stacjonarnej gospodarki Gale’a z tzw. wielopasmową magistralą, złożoną z wiązki promieni von Neumanna oraz udowodniliśmy „słabe” twierdzenie o magistrali głoszące, że w długich okresach czasu (w długim horyzoncie $T = \{0, 1, \dots, t_1\}$) optymalne procesy wzrostu w takiej gospodarce są prawie zawsze zbieżne do wielopasmowej magistrali². Obecnie prezentujemy „silną” oraz „bardzo silną” wersję twierdzenia o wielopasmowej magistrali w stacjonarnej gospodarce Gale’a. Pierwsze głosi, że zbieżność optymalnych procesów wzrostu do wielopasmowej magistrali ma miejsce zawsze, poza ewentualnie pewną skończoną liczbą jego początkowych i/lub końcowych okresów, niezależną od długości horyzontu T . Im dłuższy jest horyzont, tym dłużej optymalne procesy w fazie środkowej przebiegają w dowolnie bliskim otoczeniu wielopasmowej magistrali. Zgodnie z drugim twierdzeniem, „wejście” optymalnego procesu wzrostu na wielopasmową magistralę jest (prawie) bezpowrotne, po osiągnięciu magistrali gospodarka pozostaje na niej wszędzie dalej poza, być może, ostatnim okresem horyzontu T .

Wartością dodaną artykułu – na tle znanej autorowi literatury przedmiotu – jest wykazanie, że magistralna stabilność optymalnych procesów wzrostu jest nie tylko atrybutem stacjonarnej gospodarki Gale’a z pojedynczą magistralą, lecz także jej immanentną cechą również wtedy, gdy miejsce pojedynczej magistrali (pojedynczego promienia von Neumanna) w gospodarce zajmuje wiązka magistral (magistrala wielopasmowa).

2. MODEL

W modelu stacjonarnej gospodarki Gale’a z wielopasmową magistralą przedstawionym w pracy Panek (2016b) kluczową rolę gra przestrzeń produkcyjna typu Gale’a $Z \subset R_+^{2n}$ złożona z par n -wymiarowych, nieujemnych wektorów nakładów (zużycia) $x = (x_1, \dots, x_n)$ i produkcji (wyników) $y = (y_1, \dots, y_n)$. Zapis $(x, y) \in Z$ oznacza, że

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Ekonomii Matematycznej, al. Niepodległości 10, 60-967 Poznań, Polska, e-mail: emil.panek@ue.poznan.pl.

² W sensie miary kątowej, którą definiujemy dalej (patrz p. 2). Na wielopasmowej magistrali gospodarka osiąga najwyższe tempo wzrostu (tamże).

w gospodarce z wektora nakładów x można wytworzyć (w ustalonej jednostce czasu) wektor produkcji y . Przestrzeń produkcyjna Gale'a spełnia następujące warunki:

$$(G1) \forall (x, y) \in Z \forall \lambda \geq 0 (\lambda(x, y) \in Z)$$

(warunek proporcjonalności nakładów i wyników),

$$(G2) \forall (x^i, y^i) \in Z, i = 1, 2 ((x^1 + x^2, y^1 + y^2) \in Z)$$

(warunek addytywności procesów produkcyjnych),

$$(G3) \forall (x, y) \in Z (x = 0 \Rightarrow y = 0)$$

(warunek „braku rogu obfitości”),

$$(G4) \forall (x, y) \in Z (x' \geq x \Rightarrow (x', y) \in Z)$$

(możliwość marnotrawstwa nakładów),

$$(G5) \forall (x, y) \in Z (0 \leq y' \leq y \Rightarrow (x, y') \in Z)$$

(możliwość marnotrawstwa mocy produkcyjnych)

$$(G6) \forall (x^i, y^i) \in Z, i = 1, 2, \dots, \infty \left((x^i, y^i) \xrightarrow{i} (\bar{x}, \bar{y}) \Rightarrow (\bar{x}, \bar{y}) \in Z \right)$$

(domkniętość przestrzeni produkcyjnych).

Zbiór Z jest stożkiem domkniętym i wypukłym, z wierzchołkiem w 0. Przedmiotem naszego zainteresowania są nietrywialne (niezerowe) procesy $(x, y) \in Z \setminus \{0\}$. Liczbę $\alpha(x, y) = \max\{\alpha | \alpha x \leq y\}$ nazywamy wskaźnikiem technologicznej efektywności procesu $(x, y) \in Z \setminus \{0\}$. Można wykazać, że:

- funkcja $\alpha(\cdot)$ jest ciągła i dodatnio jednorodna stopnia 0 na $Z \setminus \{0\}$,
- istnieje $\alpha_M = \max_{(x, y) \in Z \setminus \{0\}} \alpha(x, y) = \alpha(\bar{x}, \bar{y}) < +\infty$,

zob. np. Panek (2003, tw. 5.2). Proces $(\bar{x}, \bar{y}) \in Z \setminus \{0\}$, dla którego $\alpha(\bar{x}, \bar{y}) = \alpha_M$, nazywamy optymalnym procesem produkcji. Proces ten jest określony z dokładnością do mnożenia przez stałą dodatnią (z dokładnością do struktury). Liczbę α_M nazywamy optymalnym wskaźnikiem technologicznej efektywności produkcji w gospodarce Gale'a. Przez

$$Z_{opt} = \{(\bar{x}, \bar{y}) \in Z \setminus \{0\} | \alpha(\bar{x}, \bar{y}) = \alpha_M\} \quad (1)$$

oznaczamy zbiór wszystkich optymalnych procesów produkcji w gospodarce Gale'a. Zbiór Z_{opt} jest stożkiem wypukłym nie zawierającym 0, Panek (2016b, tw. 1). Zakładamy, że gospodarka Gale'a spełnia następujący warunek silnej regularności:

$$(G7) \forall (\bar{x}, \bar{y}) \in Z_{opt} (\bar{y} > 0)$$

głoszący, że w optymalnych procesach produkcji wytwarzane są wszystkie towary³. Wynika stąd, że technologiczna efektywność każdego optymalnego procesu produkcji jest dodatnia:

$$\forall (\bar{x}, \bar{y}) \in Z_{opt} (\alpha(\bar{x}, \bar{y}) = \alpha_M > 0).$$

Jeżeli $(x, y) \in Z \setminus \{0\}$ i $y \neq 0$, to o wektorze $s = \frac{y}{\|y\|} = \left(\frac{y_1}{\|y\|}, \frac{y_2}{\|y\|}, \dots, \frac{y_n}{\|y\|} \right)$ mówimy, że charakteryzuje strukturę produkcji w procesie (x, y) ⁴. Przez

$$S = \left\{ s \mid \exists (\bar{x}, \bar{y}) \in Z_{opt} \left(s = \frac{\bar{y}}{\|\bar{y}\|} \right) \right\} \quad (2)$$

oznaczamy zbiór wektorów charakteryzujących strukturę produkcji we wszystkich optymalnych procesach $(\bar{x}, \bar{y}) \in Z_{opt}$. Przy założeniach **(G1)–(G7)** zbiór ten jest niepusty, zwarty i wypukły oraz składa się (wobec **(G7)**) wyłącznie z wektorów dodatnich, Panek (2016b, tw. 2). Niech $s \in S$. Półprostą

$$N_s = \{\lambda s \mid \lambda > 0\}$$

nazywamy promieniem von Neumanna (generowanym przez proces $(\bar{x}, \bar{y}) \in Z_{opt}$, $s = \frac{\bar{y}}{\|\bar{y}\|}$). Zbiór (wiązkę) promieni

$$\mathbb{N} = \bigcup_{s \in S} N_s$$

nazywamy wielopasmową magistralą produkcyjną w stacjonarnej gospodarce Gale'a.

Niech $p = (p_1, \dots, p_n) \geq 0$ ⁵ będzie wektorem cen towarów oraz $(x, y) \in Z \setminus \{0\}$. Liczbę⁶

$$\beta(x, y, p) = \frac{\langle p, y \rangle}{\langle p, x \rangle}$$

($\langle p, x \rangle \neq 0$) nazywamy wskaźnikiem ekonomicznej efektywności procesu (x, y) przy cenach p . Jeżeli zachodzą warunki **(G1)–(G7)**, to istnieją takie ceny $\bar{p} \geq 0$, że

$$\forall (x, y) \in Z \setminus \{0\} (\beta(x, y, \bar{p}) \leq \alpha_M)$$

³ Słabszy warunek (G7') przedstawimy w p. 4.

⁴ Tutaj i dalej: jeżeli $a \in R^n$, to $\|a\| = \sum_{i=1}^n |a_i|$. Zauważmy przy okazji, że (wobec **(G3)**) jeżeli $(x, y) \in Z \setminus \{0\}$, to $x \neq 0$.

⁵ Zapis $p \geq 0$ oznacza, że $p \geq 0$ oraz $p \neq 0$.

⁶ Tutaj i dalej: jeżeli $a, b \in R^n$, to $\langle a, b \rangle = \sum_{i=1}^n a_i b_i$.

(wszędzie gdzie $\langle p, x \rangle \neq 0$) oraz

$$\forall (\bar{x}, \bar{y}) \in Z_{opt} \setminus \{0\} (\beta(\bar{x}, \bar{y}, \bar{p}) = \alpha_M)$$

lub równoważnie

$$\forall (x, y) \in Z \setminus \{0\} (\langle \bar{p}, y \rangle \leq \alpha_M \langle \bar{p}, x \rangle) \quad (3)$$

oraz

$$\forall (\bar{x}, \bar{y}) \in Z_{opt} \setminus \{0\} (\langle \bar{p}, \bar{y} \rangle = \alpha_M \langle \bar{p}, \bar{x} \rangle > 0), \quad (4)$$

Panek (2016b, tw. 4). Każda trójka $\{\alpha_M, (\bar{x}, \bar{y}), \bar{p}\}$ wyznacza tzw. optymalny stan równowagi v. Neumanna. Wektor $\bar{p}$ nazywamy wektorem cen równowagi.

Zakładamy, że czas biegnie skokowo oraz $T = \{0, 1, \dots, t_1\}$ jest ustalonym interesującym nas horyzontem funkcjonowania gospodarki ($t_1 < +\infty$). Przez $x(t) = (x_1(t), \dots, x_n(t))$ oznaczamy wektor nakładów, a przez $y(t) = (y_1(t), \dots, y_n(t))$ wektor produkcji w okresie t . Warunek $(x(t), y(t)) \in Z$ oznacza, że w okresie t z nakładów $x(t)$ możliwe jest wytworzenie produkcji $y(t)$. Nakłady w okresie następnym $x(t+1)$ mogą pochodzić jedynie z produkcji $y(t)$ wytworzonej w okresie poprzednim, $x(t+1) \leq y(t)$. W świetle **(G4)** prowadzi to do warunku:

$$(y(t), y(t+1)) \in Z, \quad t = 0, 1, \dots, t_1. \quad (5)$$

W okresie początkowym produkcja jest ustalona:

$$y(0) = y^0 \geq 0. \quad (6)$$

O ciągu wektorów produkcji $\{y(t)\}_{t=0}^{t_1}$ spełniających warunki (5)–(6) mówimy, że opisuje (y^0, t_1) – dopuszczalny proces wzrostu w gospodarce Gale’a. Dla dowolnego początkowego wektora produkcji $y^0 \geq 0$ i dowolnej długości horyzontu T istnieją (y^0, t_1) – dopuszczalne procesy wzrostu, zob. np. Panek (2003, tw. 5.7). W gospodarce Gale’a szczególną postać mają dopuszczalne procesy $\{\bar{y}(t)\}_{t=0}^{t_1}$, w których $\bar{y}(t+1) = \alpha_M \bar{y}(t)$, lub inaczej procesy postaci

$$\bar{y}(t) = \alpha_M^t \bar{y}^0, \quad (7)$$

nazywane stacjonarnymi procesami z tempem wzrostu α_M . Przy założeniach **(G1)–(G6)** stacjonarne procesy wzrostu postaci (7) istnieją, gdy początkowy wektor produkcji spełnia warunek $\bar{y}^0 \in \mathbb{N}$ (równoważnie, gdy $\frac{\bar{y}^0}{\|\bar{y}^0\|} = s \in S$). Są one określone z dokładnością do struktury (mnożenia przez dowolną stałą dodatnią): jeżeli $\bar{y}^0 \in \mathbb{N}$, to

$$\forall t \in T (\bar{y}(t) \in \mathbb{N}) \text{ oraz } \forall t \in T \forall \lambda > 0 (\lambda \bar{y}(t) \in \mathbb{N}).$$

W gospodarce Gale'a tempo α_M jest najwyższym możliwym do osiągnięcia tempem wzrostu, w szczególności nie istnieją stacjonarne procesy wzrostu postaci $\bar{y}(t) = \gamma^t \bar{y}^0$ z tempem $\gamma > \alpha_M$. Dlatego procesy postaci (7) nazywane są także optymalnymi stacjonarnymi procesami w gospodarce Gale'a. Wszystkie takie procesy leżą na wielopasmowej magistrali $\mathbb{N}$.

Wszędzie poza magistralą osiąganą przez gospodarkę tempo wzrostu jest niższe od α_M . Z faktem tym koresponduje założenie, że również efektywność ekonomiczna każdego procesu poza magistralą $\mathbb{N}$ jest niższa od optymalnej⁷. W szczególności oznacza to, że

$$(G8) \quad \forall (x, y) \in Z \setminus \{0\} \quad (x \notin \mathbb{N} \Rightarrow \beta(x, y, \bar{p}) < \alpha_M).$$

Ustalmy następującą miarę odległości wektora towarów (w zależności od kontekstu będą to nakłady x lub produkcja y) od wielopasmowej magistrali $\mathbb{N}$ ⁸:

$$d(x, \mathbb{N}) = \inf_{x' \in \mathbb{N}} \left\| \frac{x}{\|x\|} - \frac{x'}{\|x'\|} \right\|.$$

Jeżeli zachodzi warunek (G8), to

$$\forall \varepsilon > 0 \exists \delta_\varepsilon \in (0, \alpha_M) \forall (x, y) \in Z \setminus \{0\} \quad (d(x, \mathbb{N}) \geq \varepsilon \Rightarrow \beta(x, y, \bar{p}) \leq \alpha_M - \delta_\varepsilon), \quad (8)$$

Panek (2016b, tw. 5).

Niech proces $\{y^*(t)\}_{t=0}^{t_1}$ będzie rozwiązaniem następującego zadania maksymalizacji wartości produkcji mierzonej w cenach równowagi w ostatnim okresie t_1 horyzontu T :

$$\max \langle \bar{p}, y(t_1) \rangle \quad (9)$$

$$\text{p.w. (5)–(6)}. \quad (10)$$

Przy założeniach (G1)–(G6) zadanie to ma rozwiązanie⁹. Nazywamy je $(y^0, t_1, \bar{p})$ – optymalnym procesem wzrostu w stacjonarnej gospodarce Gale'a.

⁷ W klasycznym modelu Gale'a (z pojedynczym promieniem von Neumanna) warunek ten zapewnia jednoznaczność magistrali, zob. np. Nikaido (1968, rozdz. IV, § 13), Panek (2003, rozdz. 5, pkt 5.1.3), Takayama (1985, rozdz. 7).

⁸ W ekonomii matematycznej wielkość $\left\| \frac{x}{\|x\|} - \frac{x'}{\|x'\|} \right\|$ nazywana bywa umownie odległością kątową między wektorami x , x' ; zob. np. Nikaido (1968, rozdz. 4, p. 13.3). Przez analogię $d(x, \mathbb{N})$ nazywamy odległością kątową wektora x od wielopasmowej magistrali $\mathbb{N}$.

⁹ Zob. np. Panek (2003, tw. 5.7).

□ **Twierdzenie 1** („Słabe” twierdzenie o magistrali). Jeżeli zachodzą warunki **(G1)–(G8)** oraz istnieje taki $(y^0, \check{t})$ – dopuszczalny proces $\{\check{y}(t)\}_{t=0}^{\check{t}}$, $\check{t} < t_1$, że $\check{y}(\check{t}) \in \mathbb{N}$, to dla dowolnej liczby $\varepsilon > 0$ istnieje taka liczba naturalna k , iż liczba okresów czasu, w których $(y^0, t_1, \bar{p})$ – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ spełnia warunek:

$$d(y^*(t), \mathbb{N}) \geq \varepsilon$$

nie przekracza k . Liczba k zależy od ε oraz nie zależy od długości horyzontu T .

Dowód. Panek (2016b, tw. 6). ■

3. „SILNE” I „BARDZO SILNE” TWIERDZENIE O MAGISTRALI

Weźmy liczbę $\varepsilon > 0$ i oznaczmy przez $Z(\varepsilon)$ zbiór dopuszczalnych procesów produkcji z nakładami odległymi o co najmniej ε od wielopasmowej magistrali $\mathbb{N}$:

$$Z(\varepsilon) = \{(x, y) \in Z \setminus \{0\} \mid d(x, \mathbb{N}) \geq \varepsilon\}.$$

Zgodnie z (8) efektywność ekonomiczna $\beta(x, y, \bar{p})$ każdego procesu $(x, y) \in Z(\varepsilon)$ jest niższa od optymalnej o co najmniej $\delta_\varepsilon > 0$. Poniżej przedstawiamy niektóre własności funkcji $\beta(\cdot, \cdot, \bar{p})$, które będą nam potrzebne dalej. Obowiązują warunki **(G1)–(G8)**.

□ **Fakt 1.**

$$\beta(\cdot, \cdot, \bar{p}) \in C^0(Z(\varepsilon) \rightarrow R_+^1).$$

Dowód. Funkcja $\beta(\cdot, \cdot, \bar{p})$ jest ciągła na zwartym zbiorze

$$V(\varepsilon) = \{(x, y) \in Z \mid \|x\| = 1 \wedge d(x, \mathbb{N}) \geq \varepsilon\},$$

zob. Panek (2016b, dowód tw. 5). Stąd oraz z dodatniej jednorodności stopnia 0 wynika jej ciągłość na $Z(\varepsilon) = \{\lambda \cdot V(\varepsilon) \mid \lambda > 0\}$. ■

□ **Fakt 2.**

$$\forall \varepsilon > 0 \exists \max_{(x,y) \in Z(\varepsilon)} \beta(x, y, \bar{p}) = b(\varepsilon) < \alpha_M. \quad (11)$$

Dowód. Ciągła funkcja $\beta(\cdot, \cdot, \bar{p})$ jest dodatnio jednorodna stopnia 0, więc

$$\max_{(x,y) \in Z(\varepsilon)} \beta(x, y, \bar{p}) = \max_{(x,y) \in V(\varepsilon)} \beta(x, y, \bar{p}).$$

Zbiór $V(\varepsilon)$ jest zwarty, więc istnieje

$$\max_{(x,y) \in V(\varepsilon)} \beta(x, y, \bar{p}) = b(\varepsilon).$$

Zgodnie z (8):

$$\max_{(x,y) \in Z(\varepsilon)} \beta(x, y, \bar{p}) = \max_{(x,y) \in V(\varepsilon)} \beta(x, y, \bar{p}) = b(\varepsilon) \leq \alpha_M - \delta_\varepsilon < \alpha_M.$$

■

Wprowadzając oznaczenie $\delta(\varepsilon) = \alpha_M - b(\varepsilon)$ możemy warunek (11) zapisać inaczej tak¹⁰:

$$\forall \varepsilon > 0 \exists \delta(\varepsilon) \in (0, \alpha_M] \forall (x, y) \in Z(\varepsilon) (\beta(x, y, \bar{p}) \leq \alpha_M - \delta(\varepsilon)). \quad (12)$$

Funkcja $b(\cdot)$ jest nierosnąca (tam gdzie jest określona), $b(\varepsilon) \rightarrow \alpha_M$ przy $\varepsilon \rightarrow 0$, zatem $\delta(\varepsilon) \rightarrow 0$, gdy $\varepsilon \rightarrow 0$. Przy dowodzie „silnego” twierdzenia o magistrali (twierdzenie 2) korzystamy z następującego warunku monotoniczności:

(G9) Efektywność ekonomiczna procesu $(x, y) \in Z \setminus \{0\}$ maleje w miarę oddalania się (w sensie metryki d) nakładów x od wielopasmowej magistrali $\mathbb{N}$.

Jeżeli zachodzi ten warunek, wtedy funkcja $b(\cdot)$ maleje (równoważnie, funkcja $\delta(\cdot)$ rośnie na obszarze określoności¹¹).

□ **Fakt 3.**

$$\forall s \in S_{++}^n(1) \forall \bar{s} \in S \exists \sigma(s, \bar{s}) \in (0, 1] \forall \delta \in (0, \alpha_M) \exists \varepsilon' > 0$$

$$\left(\|s - \bar{s}\| < \varepsilon' \Rightarrow ((s, \sigma(s, \bar{s})\alpha_M\bar{s}) \in \Omega(1) \wedge \alpha(s, \sigma(s, \bar{s})\alpha_M\bar{s}) \geq \alpha_M - \delta) \right),$$

gdzie: $S_{++}^n(1) = \{x \in R^n \mid x > 0 \wedge \|x\| = 1\}$, $\Omega(1) = \{(x, y) \in Z \mid \|x\| = 1\}$.

Dowód¹². Weźmy dowolne wektory $s \in S_{++}^n(1)$, $\bar{s} \in S$. Zważywszy że $\|s\| = \|\bar{s}\| = 1$ stwierdzamy, iż istnieje liczba

$$\lambda(s, \bar{s}) = \min\{\lambda \mid \lambda s \geq \bar{s}\} = \max_i \frac{\bar{s}_i}{s_i} \geq 1.$$

¹⁰ Liczba $\delta_\varepsilon = \delta(\varepsilon)$ jest największą liczbą spełniającą warunek (8).

¹¹ Warunek **(G9)** jest spełniony w szczególności, gdy wszędzie poza Z_{opt} przestrzeń produkcyjna Gale'a jest stożkiem silnie wypukłym: dla dowolnych liczb $\alpha, \beta > 0$ i każdej pary liniowo niezależnych procesów $(x^1, y^1), (x^2, y^2) \in Z \setminus Z_{opt}$ ich kombinacja liniowa $(x, y) = \alpha(x^1, y^1) + \beta(x^2, y^2)$ jest punktem wewnętrznym przestrzeni produkcyjnej Z .

¹² Dowód wzorowany na Panek (2015, lemat 1).

Ponieważ $(\bar{s}, \alpha_M \bar{s}) \in \Omega(1) \subset Z$, więc z **(G4)** dostajemy $(\lambda(s, \bar{s})s, \alpha_M \bar{s}) \in Z$, czyli

$$(s, \sigma(s, \bar{s})\alpha_M \bar{s}) \in \Omega(1), \quad (13)$$

gdzie $\sigma(s, \bar{s}) = \frac{1}{\lambda(s, \bar{s})}$, $\sigma \in C^0(S_{++}(1) \times S \rightarrow (0, 1])$, $\sigma(\bar{s}, \bar{s}) = 1$.

Funkcja α jest ciągła na $\Omega(1)$ oraz

$$\alpha(\bar{s}, \alpha_M \bar{s}) = \max_{(x, y) \in \Omega(1)} \alpha(x, y) = \alpha_M,$$

więc

$$\forall \delta \in (0, \alpha_M) \exists \varepsilon' > 0 (\|s - \bar{s}\| < \varepsilon' \Rightarrow \alpha(s, \sigma(s, \bar{s})\alpha_M \bar{s}) \geq \alpha_M - \delta). \quad (14)$$

Z (13), (14) wynika teza. ■

Fakt 3 głosi, że jeżeli wektor (unormowanych) nakładów s leży blisko któregoś wektora $\bar{s}$ (struktury produkcji na wielopasmowej magistrali $\mathbb{N}$), to istnieje proces, który (z efektywnością dowolnie bliską optymalnej) z nakładów s prowadzi do wielopasmowej magistrali $\mathbb{N}$.

□ **Twierdzenie 2** („Silne” twierdzenie o magistrali). Niech $\{y^*(t)\}_{t=0}^{t_1}$ będzie $(y^0, t_1, \bar{p})$ – optymalnym procesem wzrostu. Jeżeli spełnione są warunki **(G1)–(G9)** oraz istnieje taki $(y^0, \check{t})$ – dopuszczalny proces $\{\check{y}(t)\}_{t=0}^{\check{t}}$, $\check{t} < t_1$, że $\check{y}(\check{t}) \in \mathbb{N}$, to

$$\forall \varepsilon > 0 \exists k_\varepsilon \in \mathbb{N} (t_1 > 2k_\varepsilon \Rightarrow \forall t \in \{k_\varepsilon, k_\varepsilon + 1, \dots, t_1 - k_\varepsilon\} (d(y^*(t), \mathbb{N}) < \varepsilon))$$

(N oznacza tutaj zbiór liczb naturalnych).

Dowód. Pokażemy najpierw, że

$$\exists \tilde{\varepsilon} > 0 \forall s \in S (U_{\tilde{\varepsilon}}(s) \subset R_{++}^n), \quad (15)$$

gdzie $U_{\tilde{\varepsilon}}(s) = \{x \in R^n \mid \|x - s\| < \tilde{\varepsilon}\}$. Istotnie, ponieważ zbiór $S \subset R_{++}^n$ jest zwarty i składa się z wektorów dodatnich, to

$$\forall i \exists v_i = \min_{s \in S} s_i > 0.$$

Niech

$$f(s, y) = s + y,$$

$s \in S$, $y \in K_v = \{x \in R^n \mid -v \leq x_i \leq v, i = 1, \dots, n\}$, $v = \frac{1}{2} \min_i v_i > 0$. Zbiór $S \times K_v$ jest zwarty, a odwzorowanie $f: S \times K_v \rightarrow R^n$ ciągłe, więc jego obraz $G_v = f(S \times K_v)$ jest zwarty oraz $S \subset G_v \subset R_{++}^n$. Niech $\tilde{\varepsilon} = \frac{v}{2} > 0$. Wtedy $\forall s \in S$ ($U_{\tilde{\varepsilon}}(s) \subset G_v \subset R_{++}^n$).

Weźmy teraz dowolną liczbę $\varepsilon > 0$, liczbę $\delta(\varepsilon) \in (0, \alpha_M]$ spełniającą warunek (12), liczbę $\delta \in (0, \delta(\varepsilon))$ oraz odpowiadającą jej liczbę $\varepsilon' \in (0, \min\{\varepsilon, \tilde{\varepsilon}\})$ z faktu 3. Zgodnie ze „słabym” twierdzeniem o magistrali istnieje taka liczba naturalna $k_{\varepsilon'}$, że jeżeli $t_1 > k_{\varepsilon'}$, to

$$d(y^*(t), \mathbb{N}) < \varepsilon' \quad (16)$$

co najmniej w jednym okresie $t < t_1$. Niech $t_1 > 2k_{\varepsilon'}$. Przez τ_1 oznaczmy pierwszy, a przez τ_2 ostatni okres, w którym zachodzi warunek (16). Wówczas (w świetle (15), zważywszy że $\varepsilon' \leq \tilde{\varepsilon}$) mamy

$$\frac{y^*(\tau_1)}{\|y^*(\tau_1)\|} > 0$$

oraz w myśl faktu 3:

$$\exists \bar{s} \in \bar{S} \left(\left(\frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}, \sigma^* \alpha_M \bar{s} \right) \in \Omega(1) \right),$$

tnz.

$$(y^*(\tau_1), \rho \sigma^* \alpha_M \bar{s}) \in Z,$$

gdzie $\sigma^* = \sigma(s^*(\tau_1), \bar{s})$, $s^*(\tau_1) = \frac{y^*(\tau_1)}{\|y^*(\tau_1)\|} > 0$, $\rho = \|y^*(\tau_1)\| > 0$. Dalej dowód przebiega analogicznie jak w pracy Panek (2015, tw. 3, s. 255–256). Tworzymy (y^0, t_1) – dopuszczalny proces

$$\tilde{y}(t) = \begin{cases} y^*(t), & t = 0, 1, \dots, \tau_1, \\ \rho \sigma^* \alpha_M^{t-\tau_1} \bar{s}, & t = \tau_1 + 1, \dots, t_1 \end{cases}$$

i z definicji procesu optymalnego $\{y^*(t)\}_{t=0}^{t_1}$ dostajemy nierówność:

$$\langle \bar{p}, y^*(t_1) \rangle \geq \langle \bar{p}, \tilde{y}(t_1) \rangle = \rho \sigma^* \alpha_M^{t_1-\tau_1} \langle \bar{p}, \bar{s} \rangle > 0. \quad (17)$$

Zakładając, że k' jest liczbą okresów (między τ_1 , τ_2), w których

$$d(y^*(t), \mathbb{N}) \geq \varepsilon,$$

z (3), (5), (8) otrzymujemy:

$$\langle \bar{p}, y^*(t_1) \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1 - \tau_2} \alpha_M^{\tau_2 - t_1 - k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle. \quad (18)$$

Łącząc (17), (18) dochodzimy do nierówności

$$0 < \rho \sigma^* \alpha_M^{t_1 - \tau_1} \langle \bar{p}, \bar{s} \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1 - \tau_2} \alpha_M^{\tau_2 - t_1 - k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle$$

(z której wynika w szczególności, że $\delta(\varepsilon) < \alpha_M$), lub inaczej:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\rho \sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, y^*(\tau_1) \rangle}.$$

Nierówność powyższą, po podstawieniu $s^*(\tau_1) = \frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}$, możemy zapisać w równoważnej postaci:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle}. \quad (19)$$

Zgodnie z faktem 3:

$$\alpha = \alpha(s^*(\tau_1), \sigma^* \alpha_M \bar{s}) \geq \alpha_M - \delta,$$

zatem $\alpha s^*(\tau_1) \leq \sigma^* \alpha_M \bar{s}$ oraz $(\alpha_M - \delta) s^*(\tau_1) \leq \sigma^* \alpha_M \bar{s}$. Wówczas:

$$(\alpha_M - \delta) \langle \bar{p}, s^*(\tau_1) \rangle \leq \sigma^* \alpha_M \langle \bar{p}, \bar{s} \rangle,$$

czyli

$$\frac{\sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle} \geq \frac{\alpha_M - \delta}{\alpha_M}.$$

Stąd i z (19) dostajemy:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\alpha_M - \delta}{\alpha_M}.$$

Pamiętając, że $0 < \delta(\varepsilon') < \delta(\varepsilon) < \alpha_M$ (gdyż $0 < \varepsilon' < \varepsilon$ i funkcja $\delta(\cdot)$ jest rosnąca; zob. warunek **(G9)**) oraz $\delta \in (0, \delta(\varepsilon))$ i $t_1 - \tau_2 \geq 0$, dochodzimy do nierówności:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} > \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - t_2} \frac{\alpha_M - \delta(\varepsilon)}{\alpha_M},$$

$$\text{czyli } \left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'-1} > 1 \text{ lub (równoważnie) } (\alpha_M - \delta(\varepsilon))^{k'-1} > \alpha_M^{k'-1}.$$

Jedyną nieujemną liczbą całkowitą spełniającą ten warunek jest $k' = 0$. Tezę twierdzenia otrzymujemy przyjmując $k_\varepsilon = k_{\varepsilon'}$. ■

Przypomnijmy, że zgodnie ze „słabym” twierdzeniem o magistrali (twierdzenie 1), w długich okresach czasu (w długim horyzoncie T) optymalne procesy wzrostu są „prawie zawsze” (w prawie wszystkich okresach) zbieżne do magistrali wielopasmowej. Twierdzenie nie rozstrzyga kiedy (w których okresach) zbieżność taka ma miejsce. Udowodnione wyżej „silne” twierdzenie o magistrali wielopasmowej precyzuje, że dzieje się tak zawsze, za wyjątkiem ewentualnie początkowych i/lub końcowych okresów horyzontu T . Im horyzont jest dłuższy, tym dłużej optymalne procesy wzrostu w jego środkowym okresie przebiegają w dowolnie bliskim otoczeniu wielopasmowej magistrali N .¹³ Szczególną sytuację mamy, gdy optymalny proces w pewnym okresie $\check{t} < t_1$ dociera do magistrali. Mówi o tym kolejne twierdzenie.

□ **Twierdzenie 3** („Bardzo silne” twierdzenie o magistrali). Jeżeli spełnione są warunki **(G1)–(G8)** i $(y^0, t_1, \bar{p})$ – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ w pewnym okresie $\check{t} < t_1$ dociera do wielopasmowej magistrali

$$y^*(\check{t}) \in N,$$

to

$$\forall t \in \{\check{t} + 1, \dots, t_1 - 1\} (y^*(t) \in N).$$

Dowód. Z definicji $(y^0, t_1, \bar{p})$ – optymalnego procesu wzrostu $\{y^*(t)\}_{t=0}^{t_1}$, zgodnie z (3) mamy:

$$\langle \bar{p}, y^*(t+1) \rangle \leq \alpha_M \langle \bar{p}, y^*(t) \rangle, \quad t = 0, 1, \dots, t_1 - 1. \quad (20)$$

Jeżeli w pewnym okresie $\check{t} < t_1$ proces ten dociera do magistrali, wówczas istnieje także (y^0, t_1) – dopuszczalny proces $\{\tilde{y}(t)\}_{t=0}^{t_1}$ postaci

¹³ Czytelnika zainteresowanego klasycznymi wersjami twierdzeń o zbieżności optymalnych procesów wzrostu w gospodarce typu Gale’a do pojedynczej magistrali N_s (a nie do wiązki magistral N) odsyłamy np. do prac Nikaido (1968 rozdz. IV), Gale (1967), McKenzie (1976, 1998, 2005 rozdz. 26), Lancaster (1968, rozdz. 11), Panek (2013, 2014, 2016a). Na marginesie warto zauważyć, że zainicjowane ponad pół wieku temu badania nad „efektem magistrali” w modelach dynamiki ekonomicznej typu input – output w ostatnim okresie znajdują coraz mocniejsze ugruntowanie w teorii sterowania optymalnego i teorii gier, zob. np. Kolokoltsov, Yang (2012), Rapoport, Cartigny (2010), Zaslavski (2015).

$$\tilde{y}(t) = \begin{cases} y^*(t), & t = 0, 1, \dots, \check{t}, \\ \sigma \bar{s} \alpha_M^{t-\check{t}}, & t = \check{t} + 1, \dots, t_1, \end{cases}$$

$\sigma > 0$, $\bar{s} \in S$, co prowadzi do nierówności:

$$\langle \bar{p}, y^*(t_1) \rangle \geq \langle \bar{p}, \tilde{y}(t_1) \rangle = \sigma \alpha_M^{t_1-\check{t}} \langle \bar{p}, \bar{s} \rangle > 0. \quad (21)$$

Założmy, że $y^*(\tau) \notin \mathbb{N}$ w pewnym okresie $\tau \in \{\check{t} + 1, \dots, t_1 - 1\}$. Wówczas

$$\exists \varepsilon > 0 \ (d(y^*(\tau), \mathbb{N}) \geq \varepsilon),$$

i zgodnie z (8) istnieje taka liczba $\delta_\varepsilon \in (0, \alpha_M)$, że

$$\langle \bar{p}, y^*(\tau + 1) \rangle \leq (\alpha_M - \delta_\varepsilon) \langle \bar{p}, y^*(\tau) \rangle. \quad (22)$$

Z (20), (22) otrzymujemy nierówność

$$\langle \bar{p}, y^*(t_1) \rangle \leq \alpha_M^{t_1-\check{t}-1} (\alpha_M - \delta_\varepsilon) \langle \bar{p}, y^*(\check{t}) \rangle, \quad (23)$$

gdzie $y^*(\check{t}) = \sigma \bar{s} > 0$. Stąd oraz z (21) wynika, że

$$\sigma \alpha_M^{t_1-\check{t}-1} (\alpha_M - \delta_\varepsilon) \langle \bar{p}, \bar{s} \rangle \geq \sigma \alpha_M^{t_1-\check{t}} \langle \bar{p}, \bar{s} \rangle > 0,$$

co jest możliwe tylko gdy $\delta_\varepsilon \leq 0$. Otrzymana sprzeczność zamyka dowód. ■

4. UWAGI KOŃCOWE

Uwaga 1. „Bardzo silne” twierdzenie o magistrali pozostaje prawdziwe po zastąpieniu warunku silnej regularności **(G7)** następującym warunkiem tzw. słabej regularności:

(G7') $\forall s \in S \exists t_s < t_1 \wedge \exists (s, t_s)$ – dopuszczalny proces $\{y(t)\}_{t=0}^{t_s}$ spełniający warunek:

$$y(0) = s, y(t_s) > 0$$

głoszącym, że z każdego miejsca (punktu) wielopasmowej magistrali $\mathbb{N}$ pewien proces wzrostu w skończonym czasie $t_s < t_1$ prowadzi do dodatniego wektora produkcji $y(t_s)$. Warunek ten jest słabszy od **(G7)**, tzn. jeżeli zachodzi **(G7)**, to zachodzi także **(G7')**¹⁴.

¹⁴ Więcej o warunku **(G7')** piszmy w artykule Panek (2016b).

Uwaga 2. Rozpatrzmy zamiast zadania (9)–(10) następujący problem maksymalizacji użyteczności produkcji w końcowym okresie horyzontu T :

$$\max u(y(t_1)) \quad (24)$$

$$\text{p.w. (5)–(6).} \quad (25)$$

O funkcji użyteczności $u: R_+^n \rightarrow R_+^1$ zakładamy, że¹⁵:

(G10) (i) funkcja $u(\cdot)$ jest ciągła, wklęsła, rosnąca i dodatnio jednorodna stopnia 1,

$$(2i) \exists a > 0 \forall y \geq 0 (u(y) \leq a \langle \bar{p}, y \rangle),$$

$$(3i) \forall s \in S (u(s) = a \langle \bar{p}, s \rangle > 0).$$

W myśl **(2i)** funkcję użyteczności aproksymuje z góry forma liniowa z wektorem kierunkowym $a\bar{p}$, gdzie a jest pewną liczbą dodatnią. Warunek **(3i)** jest równoważny z następującym:

$$\forall y \in \mathbb{N} (u(y) = a \langle \bar{p}, y \rangle),$$

co oznacza, że na wielopasmowej magistrali (dla każdego wektora $y \in \mathbb{N}$) hiperpłaszczyzna $y_{n+1} = a \langle \bar{p}, y \rangle$ jest styczna do wykresu funkcji użyteczności $u(y)$.

Przyjmijmy następujące oznaczenia:

$$R_{y^0,0} = \{y^0\},$$

$$R_{y^0,1} = \{y(1) \mid (y^0, y(1)) \in Z\},$$

$$R_{y^0,t} = \{y(t) \mid (y(t-1), y(t)) \in Z\}, \quad t = 2, \dots, t_1.$$

Wówczas zadanie (24)–(25) jest równoważne z zadaniem

$$\max u(y) \quad (24')$$

$$y \in R_{y^0,t_1}, \quad (25')$$

R_{y^0,t_1} jest zbiorem wszystkich wektorów produkcji, które w okresie t_1 jest zdolna wytworzyć gospodarka z początkowym wektorem produkcji $y(0) = y^0$ (w okresie $t = 0$).

¹⁵ Zob. np. Nikaido (1968, rozdz. IV), Panek (2016b), Takayama (1985, rozdz. 7). Warunek **(G10)** jest silniejszy od podobnego warunku **(G9)** w artykule Panek (2016b).

Zadanie to ma rozwiązanie, gdyż funkcja użyteczności u jest ciągła, a zbiór R_{y^0, t_1} jest zwarty¹⁶. Tym samym istnieje rozwiązanie zadania (24)–(25). Nazywamy je (y^0, t_1, u) – optymalnym procesem wzrostu i podobnie jak rozwiązanie zadania (9)–(10) oznaczamy przez $\{y^*(t)\}_{t=0}^{t_1}$.

Przy założeniach **(G1)–(G10)** „silne” twierdzenie o magistrali (twierdzenie 2) pozostaje prawdziwe po zastąpieniu $(y^0, t_1, \bar{p})$ – optymalnego procesu wzrostu (rozwiązania zadania (9)–(10)) przez (y^0, t_1, u) – optymalny proces wzrostu (rozwiązanie zadania (24)–(25)). Istotnie, powtarzając dosłownie dowód twierdzenia 2 – po podstawieniu $u(y^*(t_1))$ zamiast $\langle \bar{p}, y^*(t_1) \rangle$ oraz $u(\bar{s})$ zamiast $\langle \bar{p}, \bar{s} \rangle$ w warunkach (17), (18) – zważywszy na **(G10)(2i)** dochodzimy do pary nierówności:

$$u(y^*(t_1)) \geq \rho \sigma^* \alpha_M^{t_1 - \tau_1} u(\bar{s}) \quad (17')$$

$$u(y^*(t_1)) \leq a \langle \bar{p}, y^*(t_1) \rangle \leq a (\alpha_M - \delta(\varepsilon'))^{t_1 - \tau_2} \alpha_M^{\tau_2 - \tau_1 - k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle. \quad (18')$$

Łącząc (17'), (18') i uwzględniając **(G10)(3i)** dochodzimy do warunku (19). Dalej dowód przebiega analogicznie jak w twierdzeniu 2.

Uwaga 3. Przy tych samych założeniach prawdziwe pozostaje także „bardzo silne” twierdzenie o magistrali (twierdzenie 3) po zastąpieniu $(y^0, t_1, \bar{p})$ – optymalnego procesu wzrostu procesem (y^0, t_1, u) – optymalnym. Zakładając, że istnieje choćby jeden okres $\tau \in \{\check{t} + 1, \dots, t_1 - 1\}$, w którym $y^*(\tau) \notin \mathbb{N}$ i powtarzając dowód twierdzenia 3 dochodzimy do konkluzji, że (y^0, t_1, u) – optymalny proces wzrostu (rozwiązanie zadania (24)–(25)) spełnia obecnie (wobec **(G10)(2i)**) – zamiast (21), (23) – warunki następujące:

$$u(y^*(t_1)) \geq \sigma \alpha_M^{t_1 - \check{t}} u(\bar{s}) > 0, \quad (21')$$

$$u(y^*(t_1)) \leq a \langle \bar{p}, y^*(t_1) \rangle \leq a \sigma \alpha_M^{t_1 - \check{t} - 1} (\alpha_M - \delta_\varepsilon) \langle \bar{p}, \bar{s} \rangle, \quad (23')$$

skąd wynika, że

$$a \sigma \alpha_M^{t_1 - \check{t} - 1} (\alpha_M - \delta_\varepsilon) \langle \bar{p}, \bar{s} \rangle \geq \sigma \alpha_M^{t_1 - \check{t}} u(\bar{s}) > 0.$$

Warunek **(G10)(3i)** stanowi jednak, że $a \langle \bar{p}, \bar{s} \rangle = u(\bar{s})$, co prowadzi do nierówności:

$$\frac{\alpha_M - \delta_\varepsilon}{\alpha_M} \geq 1.$$

Wynika z niej, że $\delta_\varepsilon \leq 0$, a to jest sprzeczne z (8). Otrzymana sprzeczność zamyka dowód.

¹⁶ Zbiory $R_{y^0, t}$ ($t = 0, 1, \dots, t_1$) są niepuste, wypukłe oraz zwarte, zob. Panek (2003, lemat 5.1).

5. ZAKOŃCZENIE

Słabą stroną rozważanego tutaj oraz w pracy Panek (2016b) modelu Gale'a z wielopasmową magistralą jest stacjonarność gospodarki. Wprawdzie w kilku publikacjach, m.in. Panek (2013, 2014, 2016a), zajmujemy się modelami różnych niestacjonarnych gospodarek Gale'a, ale podobnie jak w innych znanych pracach z tego zakresu magistrale są tam zawsze pojedynczymi promieniami (w przypadku stacjonarnym) lub pojedynczymi krzywymi (w wariancie niestacjonarnym), leżącymi w przestrzeni fazowej stanów gospodarki.

Z tej perspektywy interesujące jest szersze spojrzenie i prześledzenie własności gospodarek typu Neumanna-Gale'a-Leontiefa z wielopasmowymi magistralami, zmienną technologią oraz różnymi kryteriami wzrostu¹⁷.

LITERATURA

- Gale D., (1967), On Optimal Development in a Multi-Sector Economy, *Review of Economic Studies*, 34 (1), 1–18.
- Kolokoltsov V., Yang W., (2012), The Turnpike Theorems for Markov Games, *Dynamic Games and Applications*, 2 (3), 294–312.
- McKenzie L. W., (1976), Turnpike Theory, *Econometrica*, 44, 841–866.
- McKenzie L. W., (1998), Turnpikes, *American Economic Review*, 88, 1–14.
- McKenzie L. W., (2005), Optimal Economic Growth, Turnpike Theorems and Comparative Dynamics, w: Arrow K. J., Intriligator M. D., (red.), *Handbook of Mathematical Economics*, ed. 2, III, 26, 1281–1355.
- Lancaster K., (1968), *Mathematical Economics*, The Macmillan Company, New York.
- Nikaido H., (1968), *Convex Structures and Economic Theory*, Acad. Press, New York.
- Panek E., (2003), *Ekonomia matematyczna*, Wyd. AE, Poznań.
- Panek E., (2013), „Słaby” i „bardzo silny” efekt magistrali w niestacjonarnej gospodarce Gale'a z graniczną technologią, *Przegląd Statystyczny*, 60 (3), 291–303.
- Panek E., (2014), Niestacjonarna gospodarka Gale'a z rosnącą efektywnością produkcji na magistrali, *Przegląd Statystyczny*, 61(1), 5–14.
- Panek E., (2015), „Silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale'a. Zagadnienie wzrostu docelowego (Autopoprawka), *Przegląd Statystyczny*, 62 (2), 253–257.
- Panek E., (2016a), „Silny” efekt magistrali w modelu niestacjonarnej gospodarki z graniczną technologią, *Przegląd Statystyczny*, 63 (2), 109–121.
- Panek E., (2016b), Gospodarka Gale'a z wieloma magistralami. „Słaby” efekt magistrali, *Przegląd Statystyczny*, 63 (4), 355–374.
- Rapaport A., Cartigny P., (2010), Turnpike Theorems by a Value Function Approach, *Control, Optimization and Calculus of Variations*, 10 (1), 123–141.
- Takayama A., (1985), *Mathematical Economics*, Cambridge Univ. Press, Cambridge.
- Zaslavski A. J., (2015), *Turnpike Theory of Continuous-Time Linear Optimal Control Problems*, Springer International Publishing, Switzerland.

¹⁷ Nie tylko z kryterium maksymalizacji wartości produkcji (mierzonej w cenach von Neumanna) czy maksymalizacji społecznej funkcji użyteczności produkcji wytworzonej w końcowym okresie horyzontu T (tzw. zadania wzrostu docelowego), ale także np. z kryterium maksymalizacji łącznej wartości produkcji wytworzonej w całym horyzoncie T (odpowiednio: maksymalizacji społecznej funkcji użyteczności zdefiniowanej na wektorach produkcji we wszystkich okresach horyzontu T).

GOSPODARKA GALE'A Z WIELOMA MAGISTRALAMI.
„SILNY” I „BARDZO SILNY” EFEKT MAGISTRALI

Streszczenie

Praca nawiązuje do artykułu Panek (2016b) zawierającego dowód tzw. „słabego” twierdzenia o wielopasmowej magistrali w stacjonarnej gospodarce typu Gale’a. Obecnie prezentujemy „silną” oraz „bardzo silną” wersję twierdzenia o wielopasmowej magistrali. Pokazujemy, że mimo uogólnienia modelu – polegającego na zastąpieniu pojedynczej magistrali (promienia von Neumanna) wiązką magistral, którą nazywamy magistralą wielopasmową – nie zmieniają się wcześniej udowodnione magistralne własności optymalnych procesów wzrostu w gospodarce Gale’a.

Słowa kluczowe: stacjonarna gospodarka Gale’a, równowaga von Neumanna, wielopasmowa magistrala, „silny” i „bardzo silny” efekt magistrali

GALE’S ECONOMY WITH MULTIPLE TURNPIKES.
“STRONG” AND “VERY STRONG” TURNPIKE EFFECT

Abstract

This paper refers to the Panek (2016b) which contained proof of “weak” multiple turnpike’s theorem in the Gale’s stationary economy. Here we present “strong” and “very strong” multiple turnpike’s theorem. We show that, despite the generalization of this model – involving the replacement of a single turnpike (von Neumann’s ray) with the bundle of turnpikes, which we call multilane turnpike – the previously proven turnpike’s main properties of the optimal growth processes of the Gale’s economy do not change.

Keywords: Gale’s stationary economy, von Neumann’s equilibrium, multilane turnpike, “strong” and “very strong” turnpike effect

ANDRZEJ SOKOŁOWSKI¹, MAŁGORZATA MARKOWSKA²ITERACYJNA METODA
LINIOWEGO PORZĄDKOWANIA OBIEKTÓW WIELOCECHOWYCH³

1. WSTĘP

Porządkowanie obiektów wielocechowych jest wykorzystywane w różnych dyscyplinach naukowych, przy sporządzaniu rankingów obiektów opisanych wieloma cechami oraz w zagadnieniach konstrukcji wskaźników agregatowych.

W literaturze polskiej tematyka ta była poruszana i dyskutowana od dość dawna, począwszy od artykułu Hellwiga (1968). Jednym z najsławniejszych zastosowań miary agregatowej i rangowania obiektów wielocechowych jest Indeks Rozwoju Społecznego (HDI). Historię powstania tej miary ciekawie opisują ul Haq (2003) i Sen (2000). Metody konstrukcji wskaźników agregatowych są różne. Niekiedy dzielono je na wzorcowe i bezwzorcowe, co jest o tyle niezręczne, że w tej metodologii zawsze występuje wzorzec (jawny lub ukryty). Pewną typową procedurę postępowania zaproponowano w *Handbook on Constructing Composite Indicators* wydanym przez OECD w 2005 roku (z perspektywy polskiego dorobku w tym zakresie można tu użyć słowa „dopiero”).

Celem pracy jest przedstawienie propozycji prostej, iteracyjnej metody liniowego porządkowania obiektów wielocechowych.

2. PORZĄDKOWANIE LINIOWE
– PROCEDURA TWORZENIA WSKAŹNIKA AGREGATOWEGO

Najpopularniejsza w polskiej literaturze procedura tworzenia wskaźnika agregatowego obejmuje etapy wymienione i krótko opisane poniżej.

Ustalenie kryterium ogólnego – jest ono niemożliwe do wyrażenia przy pomocy jednej zmiennej. Przykładami takich kryteriów mogą być: poziom rozwoju gospodarczego, jakość życia, ocena zawodnika w grach zespołowych, uciążliwość wydobywania w kopalni, jakość wyższej uczelni.

¹ Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, Katedra Statystyki, ul. Rakowicka 27, 31-510 Kraków, Polska.

² Uniwersytet Ekonomiczny we Wrocławiu, Wydział Ekonomii, Zarządzania i Turystyki, Katedra Gospodarki Regionalnej, ul. Nowowiejska 3, 58-500 Jelenia Góra, Polska, autor prowadzący korespondencję – e-mail: malgorzata.markowska@ue.wroc.pl.

³ Praca wykonana w ramach grantu Narodowego Centrum Nauki: 2015/17/B/HS4/01021.

Określenie kryteriów pomocniczych (ewentualne) – jest to pewnego rodzaju „rozebranie” kryterium ogólnego na sfery (części).

Ustalenie wstępnej listy zmiennych diagnostycznych – na podstawie wiedzy i doświadczenia badawczego, opinii ekspertów i przeglądu literatury, z uwzględnieniem dostępności danych.

Ustalenie ostatecznej listy zmiennych diagnostycznych – na tym etapie zazwyczaj zmienne grupuje się w podzbiory zmiennych skorelowanych i z każdej grupy wybiera zmienną reprezentantkę. Niestety wielu badaczy w tym etapie, niesłusznie eliminuje zmienne o współczynniku zmienności mniejszym od zadanej wartości (patrz Sokołowski, Sobolewski, 2016).

Określenie charakteru zmiennych – stymulanty, destymulanty, nominanty. Większość zmiennych ma charakter oczywisty – ich określenie jest intuicyjne, zaś w przypadkach wątpliwych warto zastosować procedurę Grabińskiego (1985) wykorzystującą fakt, że stymulanty powinny być dodatnio skorelowane ze stymulantami (podobnie jest dla destymulant) zaś ujemnie z destymulantami.

Ustalenie wag – jest to decyzja mająca kluczowe znaczenie dla ostatecznego porządku obiektów. Wydaje się, że najlepszym sposobem jest ustalenie wag na podstawie merytorycznej oceny cech dokonanej przez ekspertów. Najczęściej, w praktyce, stosuje się wagi równe, co jest wyjściem w miarę rozsądnym w myśl lekarskiej zasady *primo non nocere* (po pierwsze nie szkodzić). Nietrafnym posunięciem jest wyliczanie wag z wartości cech diagnostycznych, a szczególnie jako proporcjonalnych do współczynnika zmienności. Niewielu badaczy zdaje sobie sprawę z tego, że może występować tzw. ważenie mimowolne. Jeżeli na przykład przy ustalaniu agregatowego wskaźnika poziomu życia weźmiemy cztery cechy charakteryzujące opiekę zdrowotną, a tylko jedną charakteryzującą dochody (średnia płaca) to siłą rzeczy ta pierwsza sfera ma większy wpływ na wartość wskaźnika agregatowego. Innemu aspektowi ważenia mimowolnego poświęcona jest niniejsza praca.

Doprowadzenie cech do porównywalności – możliwe jest zastosowanie trzech podejść: normalizacji do przedziału $[0;1]$, standaryzacji oraz przekształceń ilorazowych. Na etapie tym następuje zazwyczaj ujednolicenie charakteru cech, czyli przekształcenie na stymulanty. Najpopularniejsza metoda przewiduje wykorzystanie następujących wzorów:

$$x_i^* = \frac{x_i - \min_i\{x_i\}}{\max_i\{x_i\} - \min_i\{x_i\}}, \text{ dla stymulant,} \quad (1)$$

$$x_i^* = \frac{\max_i\{x_i\} - x_i}{\max_i\{x_i\} - \min_i\{x_i\}}, \text{ dla destymulant.} \quad (2)$$

Agregacja zmiennych doprowadzonych do porównywalności – jest to z reguły wzór addytywny mający postać ważonej średniej arytmetycznej:

$$W_i = \frac{s}{\sum_{j=1}^m a_j} \sum_{j=1}^m a_j x_{ij}^*, \quad (3)$$

gdzie:

i – numer obiektu,

W_i – wartość wskaźnika agregatowego dla i -tego obiektu,

j – numer cechy,

m – liczba cech,

a_j – waga j -tej cechy,

x_{ij}^* – wartość zmiennej X_j obserwowanej na i -tym obiekcie (doprowadzonej do porównywalności),

s – współczynnik skali (zazwyczaj przyjmowany jako 1 lub 100).

Efektem zastosowania zaprezentowanej procedury są wartości wskaźnika agregatowego, które można porangować oraz przedstawić na wykresie.

Jak wspomniano, najpopularniejsza procedura sprowadzania cech do porównywalności to wykorzystanie wzorów (1) i (2). W tych przekształceniach wykorzystuje się zaobserwowane najmniejsze i największe wartości zmiennych. Jeżeli występują wartości odstające lub rozkład zmiennej jest bardzo asymetryczny, to w efekcie mamy do czynienia ze sztucznym ważeniem tej zmiennej. Przy asymetrii prawostronnej zmienna sztucznie zyskuje na zaznaczeniu, a przy lewostronnej – traci.

Dla ilustracji omawianego efektu w tabeli 1 przedstawiono wartości dwóch cech wykorzystywanych w rankingu uczelni opublikowanym przez „Perspektywy” w 2015 r. Cechy te są już doprowadzone do porównywalności poprzez przekształcenie ilorazowe polegające na dzieleniu wartości dla danej uczelni przez wartość dla uczelni najlepszej (tu mamy do czynienia ze stymulantami) i pomnożeniu wyniku przez 100. Większość zaprezentowanych w tabeli wartości Preferencji pracodawców jest większa od 50, podczas gdy Uznanie międzynarodowe jest zdominowane przez dwie uczelnie, a wiele z nich ma wręcz wartości mniejsze niż 10.

Tabela 1.

Wartości dwóch cech dla czołowych uczelni w 2015 r.

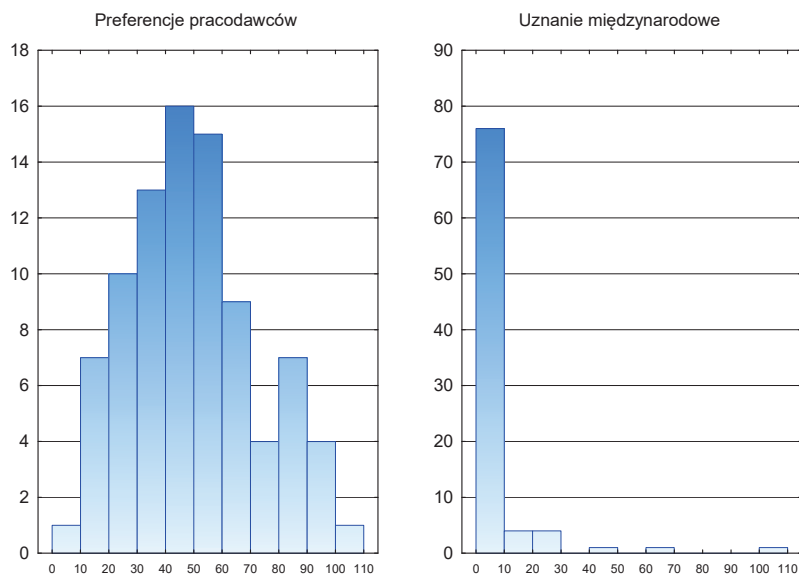
Uczelnia	Preferencje pracodawców	Uznanie międzynarodowe
Uniwersytet Jagielloński	92,62	69,43
Uniwersytet Warszawski	98,24	100,00
Uniwersytet im. Adama Mickiewicza w Poznaniu	80,82	11,57
Politechnika Warszawska	100,00	24,82
Politechnika Wrocławska	92,53	10,89
Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie	93,40	13,75
Uniwersytet Wrocławski	83,70	22,53
Warszawski Uniwersytet Medyczny	78,74	1,90
Uniwersytet Mikołaja Kopernika w Toruniu	72,20	22,53
Gdański Uniwersytet Medyczny	55,22	1,57
Politechnika Łódzka	69,25	2,93

Tabela 1. (cd.)

Uczelnia	Preferencje pracodawców	Uznanie międzynarodowe
Szkoła Główna Handlowa w Warszawie	85,15	24,41
Uniwersytet Medyczny im. Karola Marcinkowskiego w Poznaniu	44,92	1,69
Politechnika Poznańska	81,38	2,91
Uniwersytet Medyczny im. Piastów Śląskich we Wrocławiu	61,61	0,88
Uniwersytet Łódzki	65,61	16,48
Politechnika Gdańska	57,73	0,89
Uniwersytet Gdański	66,74	2,68
Uniwersytet Medyczny w Łodzi	62,52	1,18
Uniwersytet Śląski w Katowicach	65,12	4,95
Uniwersytet Ekonomiczny w Poznaniu	81,27	1,61
Politechnika Śląska w Gliwicach	84,22	4,37
Uniwersytet Medyczny w Białymstoku	46,84	1,22
Akademia Leona Koźmińskiego w Warszawie	50,61	43,17
Uniwersytet Medyczny w Lublinie	26,14	1,30

Źródło: na podstawie *Ranking szkół wyższych Perspektywy 2015* <http://www.perspektywy.pl/RSW2015/> (data dostępu 20 października 2016 r.).

Na rysunku 1 widać bardzo wyraźną asymetrię drugiej cechy i zdominowanie większości uczelni przez kilka najlepszych.



Rysunek 1. Rozkłady dwóch cech dla 100 najlepszych uczelni w rankingu „Perspektyw” z 2015 r.

Źródło: obliczenia własne.

3. PROPOZYCJA ITERACYJNEJ METODY LINIOWEGO PORZĄDKOWANIA OBIEKTÓW WIELOCECHOWYCH

Proponujemy nową, iteracyjną metodę porządkowania obiektów wielocechowych, która pozwala na ominięcie omówionej niedogodności metody klasycznej. Kolejne pozycje w rankingu wyznacza się stopniowo, po jednej w każdej iteracji, a następnie przyporządkowany obiekt jest eliminowany ze zbioru, w którym poszukujemy obiektu następnego w kolejności.

Procedurę tę można zapisać w następujących krokach:

1. Wyznaczenie minimalnych i maksymalnych wartości cech w aktualnym zbiorze (podzbiorze) obiektów.
2. Znalezienie najlepszego obiektu według procedury klasycznej (opisanej powyżej).
3. Przyporządkowanie temu obiektowi kolejnej rangi.
4. Usunięcie tego obiektu z przetwarzanego aktualnie zbioru.
5. Powrót do punktu 1.

Ten algorytm powtarzany jest aż do momentu, gdy w podzbiorze pozostaną trzy najgorsze obiekty. Im przyporządkowywane są ostatnie rangi.

Taka procedura wymaga również zaproponowania nowego sposobu wyznaczania wskaźnika agregatowego. Przy wyznaczaniu i -tej rangi wyliczamy:

$$D_{(i+1)} = \frac{w_{(i+1)}^i}{w_{(i)}^i}, \quad (4)$$

gdzie:

$D_{(i+1)}$ – współczynnik zmniejszania się wskaźnika agregatowego,

(i) – i -ta ranga,

$w_{(i)}^i$ – lokalny wskaźnik agregatowy wyliczony przy wyznaczaniu rangi i .

Ostateczna wartość wskaźnika agregatowego wyznaczana jest według wzoru (5):

$$W_{(k)} = W_{(1)}^1 \prod_{i=2}^k D_{(i)}. \quad (5)$$

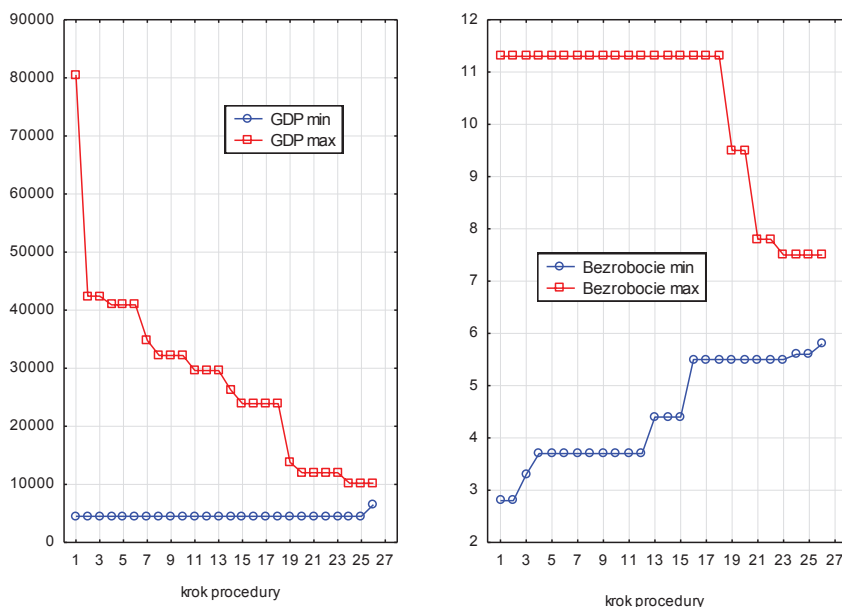
4. PRZYKŁAD ZASTOSOWANIA PROPONOWANEJ PROCEDURY

Do ilustracji zaproponowanej procedury wykorzystano dane Eurostatu dotyczące 27 krajów UE, w roku 2008. Przyjęto sześć następujących cech statystycznych:

- dochód narodowy na mieszkańca w cenach rynkowych, w euro,
- wskaźnik inflacji rocznej (w procentach),
- przeciętne dalsze trwanie życia mężczyzn w momencie urodzin, w latach,
- współczynnik zgonów niemowląt (na 1000 urodzeń żywych),
- stopa bezrobocia (w procentach),
- procent mężczyzn zagrożonych ubóstwem.

Dochód narodowy i przeciętne dalsze trwanie życia to stymulanty, zaś pozostałe cechy to destymulanty. W procedurze cechy normalizowane są według wzorów (1) i (2), ale punkty odniesienia mogą się zmieniać, gdyż obiekty rangowane są po jednym, od najlepszego do najgorszego, usuwając ze zbioru te wartości, dla których rangi zostały już wyznaczone.

Na rysunku 2 pokazano jak zmieniają się punkty odniesienia dla dwóch cech – dochodu narodowego oraz stopy bezrobocia, zaś na rysunku 3 stosunek maksimum do minimum, dla wszystkich cech, w kolejnych krokach procedury.



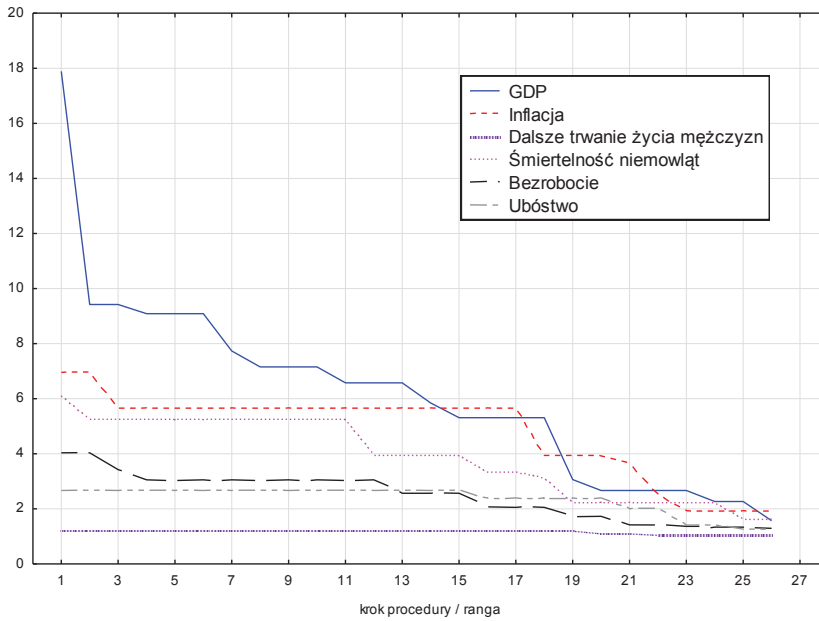
Rysunek 2. Punkty odniesienia w normalizacji zmiennych

Źródło: obliczenia własne.

Najbardziej wyraźnym zmianom ulegał stosunek największego dochodu do najmniejszego, a na zmianę punktów odniesienia największy wpływ miało wyodrębnienie (a następnie usunięcie) Luksemburga jako kraju o najwyższej randze.

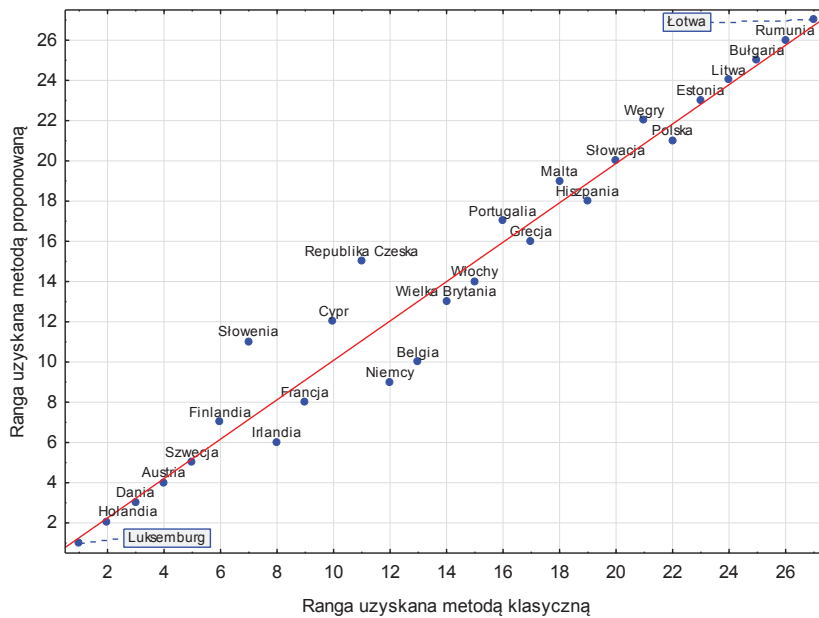
Na rysunku 4 widzimy, że przy porównaniu metody klasycznej z proponowaną, największe zmiany w porządku nastąpiły w środkowej strefie uporządkowania.

Ogólnie rzecz biorąc, dla 16 krajów z 27 analizowanych odnotowano zmiany miejsca w uporządkowaniu, w tym: poprawę o cztery pozycje dla Republiki Czeskiej oraz Słowenii, o dwie dla Cypru, a dla czterech innych krajów o jedną lokatę. Spadek o najwięcej pozycji w nowym uporządkowaniu zanotowano dla Belgii i Niemiec – o trzy, a o dwie dla Irlandii. Dla sześciu innych krajów (w tym dla Polski) – wystąpiła zmiana o jedno miejsce (por. tabela 2).



Rysunek 3. Stosunek maksimum do minimum w kolejnych krokach procedury

Źródło: obliczenia własne.



Rysunek 4. Porównanie rang uzyskanych metodą klasyczną i proponowaną metodą

Źródło: obliczenia własne.

Tabela 2.

Porównanie wyników metod klasycznej i proponowanej

Państwo	Klasyczny wskaźnik agregatowy	Proponowany wskaźnik agregatowy	Rangi według wskaźnika klasycznego	Rangi według wskaźnika proponowanego	Różnica rang
Belgia	67,5	64,3	13	10	-3
Bułgaria	28,5	20,1	25	25	0
Czechy	69,1	58,7	11	15	4
Dania	79,4	77,0	3	3	0
Niemcy	68,9	64,9	12	9	-3
Estonia	40,4	34,3	23	23	0
Irlandia	71,9	70,0	8	6	-2
Grecja	57,3	52,0	17	16	-1
Hiszpania	52,9	47,9	19	18	-1
Francja	69,5	65,3	9	8	-1
Włochy	64,0	59,6	15	14	-1
Cypr	69,2	62,4	10	12	2
Łotwa	23,1	18,1	27	27	0
Litwa	36,1	27,5	24	24	0
Luksemburg	86,3	86,3	1	1	0
Węgry	48,4	35,5	21	22	1
Malta	53,9	44,5	18	19	1
Holandia	84,5	80,3	2	2	0
Austria	79,1	75,1	4	4	0
Polska	47,1	36,1	22	21	-1
Portugalia	60,3	51,7	16	17	1
Rumunia	26,3	18,8	26	26	0
Słowenia	72,3	63,7	7	11	4
Słowacja	51,2	37,0	20	20	0
Finlandia	72,4	68,9	6	7	1
Szwecja	78,6	74,8	5	5	0
Wielka Brytania	64,7	61,3	14	13	-1

Źródło: obliczenia własne.

5. PODSUMOWANIE

Zaproponowana metoda pozwala na zredukowanie efektu tzw. ważenia mimowolnego. Procedura ustalania porządku obiektów wielocechowych ma charakter iteracyjny i pozwala na ustalenie kolejnego obiektu najlepszego spośród tych, którym jeszcze nie nadano rangi. Każdy obiekt porównywany jest tylko do tych, którym jeszcze nie zostały przyporządkowane wyższe pozycje w klasyfikacji obiektów wielowymiarowych. Ponieważ punkty odniesienia (które można traktować jako współrzędne wzorca) zmieniają się w trakcie ustalania hierarchii obiektu, zatem zaproponowano nowy sposób ustalania wartości wskaźnika agregatowego.

LITERATURA

- Grabiński T., (1985), Metody określania charakteru zmiennych w wielowymiarowej analizie porównawczej, *Zeszyty Naukowe Akademii Ekonomicznej w Krakowie*, 213, 35–63.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju oraz zasoby i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), 307–326.
- Nardo M., Saisana M., Saltelli A., Tarantola S., Hoffman A., Giovannini E., (2005), *Handbook on Constructing Composite Indicators*, OECD Statistics Working Papers, 2005/03.
- Ranking szkół wyższych Perspektywy 2015*, <http://www.perspektywy.pl/RSW2015/>.
- Sen A., (2000), A Decade of Human Development, *Journal of Human Development*, 1 (1), 17–23.
- Sokołowski A., Sobolewski M., (2016), O niewłaściwym wykorzystaniu współczynnika zmienności, XXXV Konferencja Naukowa Multivariate Statistical Analysis (MSA 2016), książka streszczeń.
- ul Haq M., (2003), The Birth of the Human Development Index, w: Fukuda-Parr S., Shiva Kuma A. K., (red.), *Readings in Human Development*, Oxford University Press, Oxford, 127–137.

ITERACYJNA METODA
LINIOWEGO PORZĄDKOWANIA OBIEKTÓW WIELOCECHOWYCH

Streszczenie

W porządkowaniu liniowym obiektów wielocechowych najpopularniejszą metodą doprowadzania danych do porównywalności jest transformacja zmiennych do przedziału $[0;1]$ z jednoczesną zamianą destymulant na stymulanty. W tym przekształceniu wykorzystuje się zaobserwowane najmniejsze i największe wartości zmiennych. Jeżeli występują wartości odstające, lub rozkład zmiennej jest bardzo asymetryczny to w efekcie mamy do czynienia ze sztucznym ważeniem tej zmiennej. Przy asymetrii prawostronnej zmienna sztucznie zyskuje na zaznaczeniu, a przy lewostronnej – traci. W pracy zaproponowano nową, iteracyjną metodę porządkowania obiektów wielocechowych, która pozwala na ominięcie omówionej niedogodności metody klasycznej. Dalsze pozycje w rankingu wyznacza się kolejno, po jednej w każdej iteracji, a przyporządkowany obiekt jest eliminowany ze zbioru, w którym poszukujemy obiektu następnego w kolejności. Taka procedura wymagała również zaproponowania nowego sposobu wyznaczania wskaźnika agregatowego.

Słowa kluczowe: porządkowanie liniowe, wskaźniki agregatowe

ITERATIVE METHOD FOR LINEAR ORDERING OF MULTIDIMENSIONAL OBJECTS

Abstract

Transformation into $[0;1]$ interval combined with changing destimulants into stimulants is the most popular method for standardizing variables used in linear ordering of multidimensional objects. Minimum and maximum values are used as reference points. Outliers and very skewed distributions result in artificial weighing of variables, imposing higher relative importance of the variable with positive skewness and lower with negative skewness. A new linear ordering procedure is proposed to overcome this problem. Ranking positions are identified one at a time, and assigned object is then eliminated from the lot. Such iterative procedure needs a new method for obtaining composite indicators, and it is also proposed in the paper.

Keywords: linear ordering, composite indicators

KAROL KUKUŁA¹, LIDIA LUTY²JESZCZE O PROCEDURZE WYBORU METODY
PORZĄDKOWANIA LINIOWEGO

1. WPROWADZENIE

Liczni autorzy budujący rankingi obiektów na podstawie oceny zjawisk złożonych stosowali dowolnie wybraną przez siebie metodę porządkowania liniowego bądź też używali do tego celu kilka metod również subiektywnie wytypowanych, poddając otrzymane wyniki opinii czytelnika. Częstość uzyskiwane wyniki przy zastosowaniu kilku metod porządkowania liniowego różnią się między sobą. W niektórych zestawieniach porównywanych par metod różnice te są bardzo wyraźne. Powstaje w związku z tym dylemat, którą z metod porządkowania wybrać.

Problematyka wyboru metody porządkowania liniowego została podjęta w pracy Kukuła, Luty (2015a). W artykule tym przedstawiono propozycję procedury wspomagającej wybór metody porządkowania liniowego. Procedura ta bazuje na mierze podobieństwa rankingów uzyskanych w wyniku zastosowania kilku metod porządkowania liniowego. Zdaniem autorów, przy wyborze metody porządkowania liniowego w badaniach, należy wziąć pod uwagę dwa postulaty:

- zrealizować wybór dostosowując metodę pod kątem wykorzystania jej własności, charakteru zmiennych diagnostycznych oraz skali pomiaru zmiennych,
- wybrać metodę, która daje najbliższe wyniki końcowe względem pozostałych metod.

Problematykę tę podjęła także Kisielińska (2016). W swej pracy przedstawiła propozycję budowy rankingu wykorzystując wskaźnik zagregowanej pozycji w rankingach otrzymanych w wyniku zastosowania kilku metod porządkowania liniowego.

Celem niniejszego artykułu jest polemika z procedurą przedstawioną w pracy Kisielińskiej (2016). Zdaniem autorki jej metoda jest mniej odporna na „rankingi odstające”, w porównaniu do procedury zaproponowanej przez Kukuła, Luty (2015a). W szczególności pod rozważę pragniemy poddać zasadność uwzględniania wszystkich rankingów, a zwłaszcza tych „odstających”, co autorka podkreśla jako atut zapropo-

¹ Uniwersytet Rolniczy, Wydział Rolniczo-Ekonomiczny, Katedra Statystyki i Ekonometrii, Al. Mickiewicza 21; 31-120 Kraków, Polska.

² Uniwersytet Rolniczy, Wydział Rolniczo-Ekonomiczny, Katedra Statystyki i Ekonometrii, Al. Mickiewicza 21; 31-120 Kraków, Polska, autor prowadzący korespondencję – e-mail: rrduka@cyf-kr.edu.pl.

nowanej procedury („wszystkie opracowane rankingi mają wpływ na pozycję obiektu w rankingu ostatecznym”). Tym samym, proponujemy procedurę eliminacji „odstających” rankingów, wskazując na przykładach zasady jej przeprowadzenia, przed sporządzeniem ostatecznego rankingu, w szczególności gdy rozważamy wykorzystanie procedury proponowanej przez Kisielińską.

2. PREZENTACJA WYBRANYCH PROCEDUR

Procedury porządkowania liniowego obiektów zostały szczegółowo omówione w literaturze przedmiotu, w tym m.in. przez takich autorów jak: Perkal (1953), Hellwig (1968), Wesołowski (1975), Bartosiewicz (1976), Nowak (1977), Strahl (1978), Borys (1978), Grabiński (1992), Kukula (2000), Lira i inni (2002), Młodak (2006), Pawełek (2008), Panek (2009), Walesiak (2014) i innych. Wielu autorów podkreśla, że rankingi tworzone różnymi metodami są niestabilne, co skłania do zwrócenia uwagi, po raz kolejny na potrzebę wypracowania procedur, które pozwolą na wybór metody lub „może”, co zaproponowała Kisielińska, procedury pozwalającej wyznaczyć ostateczne pozycje obiektów w rankingu.

Rozważmy v rankingów n obiektów opracowanych z wykorzystaniem różnych metod porządkowania liniowego opartych na zmiennej syntetycznej ze względu na stan zjawiska złożonego opisanego przez m zmiennych diagnostycznych.

W literaturze przedmiotu pojawiły się następujące procedury³:

A – wybór metody porządkowania liniowego (Kukula, Luty, 2015a), tym samym rankingu na niej podstawie opracowanego, dla której $\bar{u}_p$ jest maksymalne, gdy:

$$\bar{u}_p = \frac{1}{v-1} \sum_{\substack{q=1 \\ p \neq q}}^v m_{pq}, \quad p, q = 1, 2, \dots, v, \quad (1)$$

gdzie m_{pq} miara podobieństwa rankingów p i q (Kukula, 1989):

$$m_{pq} = 1 - \frac{2 \sum_{i=1}^n |c_{ip} - c_{iq}|}{n^2 - z}, \quad (2)$$

c_{ip} , c_{iq} – pozycja i -tego obiektu odpowiednio w rankingu p i q ;

$z = \begin{cases} 0, & n \in P \\ 1, & n \notin P \end{cases}$, a P – zbiór liczb naturalnych parzystych;

$m_{pq} \in [0, 1]$.

³ W dalszej części pracy procedury będą odpowiednio oznaczone: A lub B.

B – wyznaczenie rankingu ostatecznego (Kisielińska, 2016) na podstawie wskaźnika syntetycznego (w_i) wykorzystującego wszystkie opracowane rankingi, który definiujemy:

$$w_i = \frac{1}{v} \sum_{j=1}^v c_{ij}, \quad (3)$$

c_{ij} – pozycja i -tego obiektu w rankingu o numerze j ;

w_i – wskaźnik syntetyczny i -tego obiektu.

Sporządzony ostatecznie ranking z wykorzystaniem proponowanych procedur zależy od wyjściowego zestawu metod porządkowania. W szczególności procedura B oparta na średniej arytmetycznej rang jest czuła na „odstające” rankingi. Zasadnym zatem byłoby wyeliminowanie tych układów porządkowych. Proponujemy następujący algorytm eliminacji:

1. porównanie układów porządkowych każdy z każdym z wykorzystaniem miary m_{pq} (2) i przedstawienie ich w macierz M :

$$M = [m_{pq}] = \begin{bmatrix} 1 & m_{12} & m_{13} & \cdots & m_{1v} \\ & 1 & m_{23} & \cdots & m_{2v} \\ & & \vdots & \vdots & \vdots \\ & & & \vdots & \vdots \\ & & & & 1 \end{bmatrix}; \quad (4)$$

2. wyznaczenie średniej miary podobieństwa rankingów:

$$\bar{m} = \frac{2 \sum_{p=1}^v \sum_{p>q} m_{pq}}{v \cdot (v-1)}, \quad p, q = 1, 2, \dots, v; \quad (5)$$

3. wyznaczenie odchylenia przeciętnego:

$$d(m) = \frac{2 \sum_{p=1}^v \sum_{p>q} |m_{pq} - \bar{m}|}{v \cdot (v-1)}, \quad p, q = 1, 2, \dots, v; \quad (6)$$

4. wyeliminowanie rankingów, dla których:

$$\forall p, q, p \neq q : m_{pq} \leq \bar{m} + d(m). \quad (7)$$

3. WERYFIKACJA EMPIRYCZNA

Do sporządzenia rankingów⁴ wykorzystano następujące metody porządkowania liniowego oparte na zmiennej syntetycznej⁵:

- metody wzorcowe:
 - R1 – metoda Hellwiga, normalizacja cech z wykorzystaniem standaryzacji zmiennych,
 - R2 – metoda TOPSIS, normalizacja cech z wykorzystaniem standaryzacji zmiennych,
 - R3 – metoda pozycyjna, normalizacja cech z wykorzystaniem standaryzacji pozycyjnej z medianą Webera, agregacja zmiennych dokonana z wykorzystaniem propozycji Hellwiga;
- metody bezwzorcowe⁶, uwzględniające normalizację cech z wykorzystaniem:
 - R4 – standaryzacji zmiennych,
 - R5 – metody unitaryzacji zerowanej,
 - R6 – metody Strahl,
 - R7 – metody Nowaka.

Przyjęto ten sam zestaw metod co stosowała w swoim artykule Kisielińska (2016) celem dokonania porównań.

Przykład 1⁷

Do oceny potencjalnych możliwości zaspokojenia zapotrzebowania na produkty rolnicze państw Unii Europejskie (UE) autorka wytypowała następujące zmienne diagnostyczne:

- X_1 – liczba ludności na 1 zatrudnionego w rolnictwie,
- X_2 – liczba osób zatrudnionych w rolnictwie otrzymujących wynagrodzenie w stosunku do nie otrzymujących wynagrodzenia,
- X_3 – powierzchnia UR przypadająca na 1 zatrudnionego w rolnictwie,
- X_4 – udział UR w powierzchni całkowitej,
- X_5 – produkcja mięsa na 1 mieszkańca,
- X_6 – produkcja roślinna na 1 mieszkańca,
- X_7 – wartość dodana wypracowana przez 1 zatrudnionego w rolnictwie.

Dla wszystkich zmiennych diagnostycznych w badanej grupie obiektów współczynnik zmienności jest większy od 10%, a ilorazy skrajnych wartości przekraczają wartość dwa (tabela 1).

⁴ W dalszej części pracy metody (lub rankingi sporządzone z ich wykorzystaniem) będą oznaczone odpowiednio: R1, R2, R2, R3, R4, R5, R6, R7.

⁵ Wymienione procedury wykorzystano i omówiono m.in. w pracy Kukula, Luty (2015a).

⁶ Zmienna syntetyczna wyznaczona jest jako średnia arytmetyczna sum znormalizowanych wartości zmiennych diagnostycznych.

⁷ Przykład pochodzi z artykułu Kisielińska (2016).

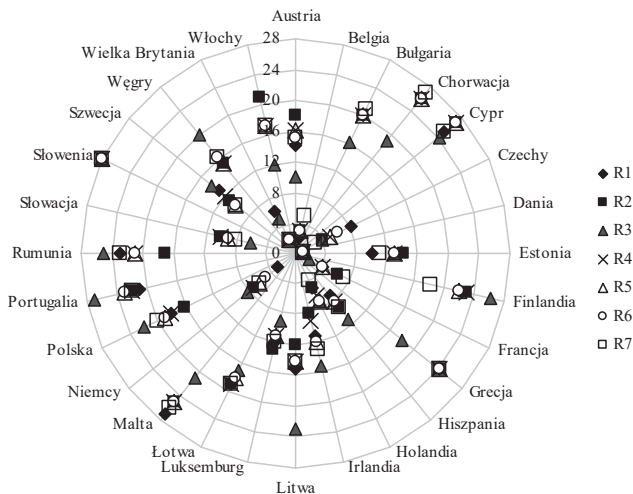
Tabela 1.

Podstawowe charakterystyki przyjętych zmiennych diagnostycznych

Charakterystyki liczbowe	Zmienne diagnostyczne						
	X_1	X_2	X_3	X_4	X_5	X_6	X_7
Maksymalna wartość	219,00	2,77	57,70	0,71	327,80	2,86	67,70
Minimalna wartość	15,00	0,06	2,16	0,07	7,20	0,16	3,99
Średnia arytmetyczna	75,36	0,57	23,85	0,40	82,15	1,20	19,61
Mediana	60,50	0,39	22,83	0,43	64,95	1,10	13,03
Współrzędne mediany Webera	60,75	0,51	21,05	0,34	69,82	1,11	16,61
Odchylenie standardowe	57,76	0,61	15,07	0,16	65,35	0,67	15,34
Współczynnik zmienności	0,77	1,07	0,63	0,39	0,80	0,56	0,78
Iloraz skrajnych wartości	14,60	46,17	26,71	10,14	45,53	17,88	16,97
Współczynnik skośności	1,02	2,34	0,52	-0,33	2,28	0,79	1,36

Źródło: obliczenia własne na podstawie z artykułu Kisielińskiej (2016).

Wykorzystując dane z baz EUROSTATU z 2015 roku sporządzono rankingi: R1, R2, ..., R7 państw UE ze względu na ocenę potencjału możliwości zaspokojenia zapotrzebowania na produkty rolnicze. Pozycje państw w sporządzonych rankingach przedstawiono na rysunku 1. Należy zwrócić uwagę, że tylko dwa spośród 28 państw UE we wszystkich rankingach miały tę samą pozycję, były to państwa sklasyfikowane odpowiednio najwyżej (Dania) i najniżej (Słowenia). Litwa w rankingu sporządzonym z wykorzystaniem metody TOPSIS (R2) uplasowała się na 12 pozycji, a gdy bazowano na metodzie pozycyjnej (R3) na pozycji 23.



Rysunek 1. Rangi państw UE ze względu na ocenę potencjału możliwości zaspokojenia zapotrzebowania na produkty rolnicze uzyskane w wyniku zastosowania siedmiu metod porządkowania liniowego obiektów (R1, R2, ..., R7)

Źródło: opracowanie własne na podstawie danych z artykułu Kisielińska (2016).

Dla każdej pary sporządzonych układów porządkowych: R1, R2, ..., R7 wyznaczono miarę m_{pq} i przedstawiono w macierzy M^8 :

$$M = \begin{bmatrix} 1,000 & 0,862 & 0,755 & 0,913 & 0,908 & 0,918 & 0,898 \\ & 1,000 & 0,714 & 0,939 & 0,913 & 0,903 & 0,878 \\ & & 1,000 & 0,765 & 0,776 & 0,776 & 0,750 \\ & & & 1,000 & 0,969 & 0,959 & 0,903 \\ & & & & 1,000 & 0,990 & 0,918 \\ & & & & & 1,000 & 0,918 \\ & & & & & & 1,000 \end{bmatrix}.$$

Następnie według wzorów (5) i (6) wyznaczono odpowiednio średnią miarę podobieństwa oraz odchylenie przeciętne: $\bar{m} = 0,873$ oraz $d(m) = 0,068$. Największe podobieństwo (macierz M) charakteryzuje parę rankingów uzyskanych z wykorzystaniem dwóch bezwzorcowych metod porządkowania, w których formułą normującą była odpowiednio metoda Strahl i metoda unitaryzacji zerowanej ($m_{56} = 0,990$). Z kolei największe zróżnicowanie charakteryzuje parę rankingów uzyskanych z wykorzystaniem metody TOPSIS i metody pozycyjnej ($m_{23} = 0,714$). Zauważymy także, że w macierzy M wszystkie wartości (oczywiście poza elementem na przekątnej) w trzecim wierszu i w trzeciej kolumnie są mniejsze od przeciętnej miary podobieństwa ($\bar{m}$), co pozwala stwierdzić, że ranking zbudowany z wykorzystaniem metody pozycyjnej (R3) jest najbardziej „odstający”.

Wykorzystując algorytm zaproponowanej eliminacji, dla danych z przykładu 1 „odstającymi” rankingami okazały się te, w których stosowano: metodę Hellwiga, metodę TOPSIS, metodę pozycyjną, metodę Nowaka. Tym samym do proponowanego wskazania/wyznaczenia ostatecznego rankingu proponuje się wykorzystać zredukowaną liczbę rankingów, zostawiając te, które uzyskano z wykorzystaniem metod bezwzorcowych, gdy do normalizacji przyjęto odpowiednio: standaryzację zmiennych (R4), metodę unitaryzacji zerowanej (R5), metodę Strahl (R6). W tabeli 2 przedstawiono rangi państw uzyskane w wyniku zastosowania procedury Kukula, Luty (A) i procedury Kisieleńska (B), wykorzystujące odpowiednio wszystkie wyjściowe rankingi (oznaczono odpowiednio: rA i rB) oraz zredukowaną liczbę rankingów (oznaczono odpowiednio: r*A i r*B).

⁸ Numer wiersza (kolumny) odpowiada numerowi w przyjętym oznaczeniu metod.

Tabela 2.

Pozycje państw UE uzyskane w wyniku zastosowania procedury Kukuła, Luty (A) i procedury Kisielińskiej (B), gdy wykorzystano odpowiednio wszystkie wyjściowe rankingi (r_A , r_B) i zredukowaną ich liczbę (r^*_A , r^*_B)

Kraj	Procedura A		Procedura B	
	r_A	r^*_A	r_B	r^*_B
Austria	16	16	14	16
Belgia	3	3	3	3
Bułgaria	20	20	20	20
Chorwacja	26	26	26	26
Cypr	27	27	27	27
Czechy	5	5	5	5
Dania	1	1	1	1
Estonia	13	13	13	13
Finlandia	22	22	22	22
Francja	4	4	4	4
Grecja	24	24	24	24
Hiszpania	8	8	8	8
Holandia	7	7	6	7
Irlandia	12	12	11	11
Litwa	14	14	15	14
Luksemburg	11	11	10	12
Łotwa	18	18	18	18
Malta	25	25	25	25
Niemcy	6	6	7	6
Polska	19	19	19	19
Portugalia	23	23	23	23
Rumunia	21	21	21	21
Słowacja	9	9	9	9
Słowenia	28	28	28	28
Szwecja	10	10	12	10
Węgry	15	15	16	15
Wielka Brytania	2	2	2	2
Włochy	17	17	17	17

Źródło: opracowanie własne na podstawie danych z artykułu Kisielińska (2016).

Zastosowanie procedury Kukula, Luty zarówno w odniesieniu do wszystkich wyjściowych rankingów jak i do zredukowanej liczby nie zmieniło końcowego uporządkowania państw. W rozpatrywanym przykładzie, w obu przypadkach najbardziej podobnym do wszystkich uwzględnianych rankingów okazał się ranking oparty na metodzie unitaryzacji zerowanej (R5). W tabeli 3 przedstawiono wartości średnich miar podobieństwa oszacowane ($\bar{u}_p$) według wzoru (1) dla każdej zastosowanej metody porządkowania liniowego.

Tabela 3.

Wartości $\bar{u}_p$ (p -numer metody) z uwzględnieniem odpowiednio: wszystkich wyjściowych rankingów i wybranych rankingów (po odrzuceniu rankingów „odstających”)

Uwzględnione rankingi	Metoda konstrukcji zmiennej syntetycznej						
	R1	R2	R3	R4	R5	R6	R7
R1, R2, ..., R7	0,876	0,868	0,760	0,908	0,912	0,911	0,878
R4, R5, R6	-	-	-	0,964	0,980	0,974	-

Źródło: opracowanie własne na podstawie danych z artykułu Kisielińska (2016).

Trudno to samo stwierdzić o wynikach zastosowania procedury Kisielińskiej. Pozycje siedmiu państw różniły się w zależności od tego czy bazowano na wszystkich wyjściowych rankingach czy na zredukowanej ich liczbie.

Podkreślić należy także, iż układy porządkowe państw UE skonstruowane z wykorzystaniem odpowiednio procedury Kukula, Luty i procedury Kisielińskiej, gdy bazowano na wszystkich rankingach różniły się dla siedmiu badanych obiektów. Miara podobieństwa tych rankingów oszacowana wg wzoru (2) wynosi 0,974. Różnice te zmniejszyły się, gdy dokonano eliminacji rankingów „odstających”.

Jako ostateczny ranking proponuje się uznać ranking r^*A lub r^*B , czyli rankingi wykorzystujące omówione procedury po dokonaniu redukcji rankingów odstających. Pozycje tylko dla dwóch państw (Irlandii, Luksemburga), w tych rankingach różnią się.

Przykład 2

Wygenerowano realizację siedmiu zmiennych diagnostycznych: $X_1, X_2, \dots, X_7$ dla siedemnastu umownych obiektów: $O1, O2, \dots, O17$, których realizację przedstawiono w tabeli 4.

Wartości charakterystyk liczbowych zmiennych diagnostycznych przedstawiono w tabeli 5. Dla wszystkich zmiennych diagnostycznych w badanej grupie obiektów współczynnik zmienności jest większy od 10%, a ilorazy skrajnych wartości przekraczają wartość dwa.

W opracowaniu przyjęto założenie, że każda zmienna jest stymulantą oraz wnosi taką samą porcję informacji do oceny badanych obiektów (wagi wszystkich zmiennych są takie same i wynoszą jeden).

Tabela 4.

Wartości zmiennych diagnostycznych

Obiekt	Zmienne diagnostyczne						
	X_1	X_2	X_3	X_4	X_5	X_6	X_7
O1	1000	102	49	1200	80	436	660
O2	250	135	42	300	75	380	520
O3	350	141	35	256	20	354	610
O4	282	155	77	340	65	440	577
O5	14	133	63	700	52	352	601
O6	268	112	45	650	41	12	502
O7	320	120	31	350	98	2	560
O8	310	135	45	280	30	442	525
O9	330	114	80	220	22	340	15
O10	260	108	48	236	66	880	10
O11	285	122	78	288	59	441	502
O12	300	5	32	688	50	350	605
O13	339	145	38	244	25	402	648
O14	340	107	36	200	100	365	504
O15	290	105	30	670	32	348	625
O16	332	144	37	210	95	400	500
O17	255	100	72	660	96	882	20

Źródło: opracowanie własne.

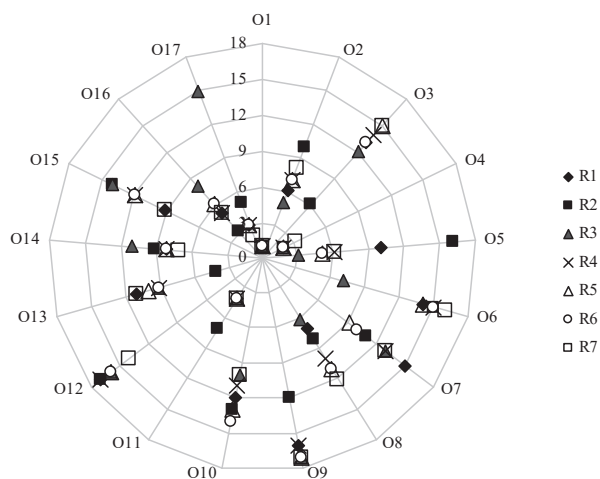
Tabela 5.

Podstawowe charakterystyki przyjętych zmiennych diagnostycznych

Charakterystyki liczbowe	Zmienne diagnostyczne						
	X_1	X_2	X_3	X_4	X_5	X_6	X_7
Maksymalna wartość	1000,00	155,00	80,00	1200,00	100,00	882,00	660,00
Minimalna wartość	14,00	5,00	30,00	200,00	20,00	2,00	10,00
Średnia arytmetyczna	325,00	116,65	49,29	440,71	59,18	401,53	469,65
Mediana	300,00	120,00	45,00	300,00	59,00	380,00	525,00
Współrzędne mediany Webera	299,50	129,05	51,16	325,69	58,39	399,13	522,24
Odchylenie standardowe	184,23	32,44	17,14	266,93	27,35	215,96	216,54
Współczynnik zmienności	0,57	0,28	0,35	0,61	0,46	0,54	0,46
Iloraz skrajnych wartości	71,43	31,00	2,67	6,00	5,00	441,00	66,00
Współczynnik skośności	2,86	-2,41	0,78	1,46	0,10	0,62	-1,63

Źródło: obliczenia własne na podstawie tabeli 4.

Wykorzystując dane z tabeli 4 sporządzono rankingi: R1, R2, ..., R7 obiektów: O1, O2, ..., O17, których pozycje są ukazane na rysunku 2. Tylko jeden obiekt (O1) we wszystkich rankingach miał stałą pozycję, a dla dwóch obiektów różnica rang wynosiła aż 13.



Rysunek 2. Rangi obiektów (O1, O2, ..., O17) uzyskane w wyniku zastosowania siedmiu metod porządkowania liniowego obiektów (R1, R2, ..., R7)

Źródło: opracowanie własne na podstawie danych z tabeli 4.

Dla każdej pary układów porządkowych: R1, R2, ..., R7 obiektów: O1, O2, ..., O17 wyznaczono wartość miary m_{pq} według wzoru (2). Wszystkie wyliczone wartości m_{pq} przedstawiono w macierzy M wg (4):

$$M = \begin{bmatrix} 1,000 & 0,667 & 0,667 & 0,875 & 0,819 & 0,819 & 0,819 \\ & 1,000 & 0,514 & 0,681 & 0,639 & 0,667 & 0,569 \\ & & 1,000 & 0,708 & 0,681 & 0,694 & 0,625 \\ & & & 1,000 & 0,903 & 0,917 & 0,875 \\ & & & & 1,000 & 0,958 & 0,847 \\ & & & & & 1,000 & 0,833 \\ & & & & & & 1,000 \end{bmatrix}.$$

Następnie korzystając ze wzorów (5) oraz (6) wyznaczono: $\bar{m} = 0,751$ oraz $d(m) = 0,110$.

Także i w tym przykładzie, największe podobieństwo (macierz M) charakteryzuje parę rankingów uzyskanych z wykorzystaniem metody Strahl i metody unitaryzacji zerowanej ($m_{56} = 0,958$). Najmniej podobną okazała się również para rankingów

uzyskanych z wykorzystaniem metody TOPSIS i metody pozycyjnej ($m_{23} = 0,514$). Zauważymy, że i w tym przykładzie wartości miar podobieństwa rankingu R3 z rankingami wyznaczonymi w oparciu o pozostałe uwzględnione metody są mniejsze od przeciętnej miary podobieństwa.

Wykorzystując algorytm zaproponowanej eliminacji „odstających” rankingów do wskazania według proponowanej procedury ostatecznego rankingu wykorzystano następujące rankingi: R1, R4, R5, R6, R7, eliminując, w tym przykładzie rankingi: R2 i R3.

Analogicznie jak w tabeli 3 dla przykładu 1 tak dla przykładu 2, w tabeli 6 przedstawiono rangi obiektów uzyskane w wyniku zastosowania dwóch procedur: A, B, wykorzystujące odpowiednio wszystkie wyjściowe rankingi (r_A , r_B) oraz zredukowaną liczbę rankingów (r^*A , r^*B).

Zastosowanie procedury A także i w tym przykładzie, w odniesieniu do wszystkich wyjściowych rankingów jak i do zredukowanej ich liczby nie zmieniło uporządkowania obiektów. Nie można tego samego stwierdzić w odniesieniu do wyników uzyskanych z wykorzystaniem procedury B. Rangi w układach porządkowych oznaczonych odpowiednio: r_B i r^*B różniły się aż dla jedenastu obiektów. Należy także zaznaczyć, że pozycje siedmiu obiektów zależały od zastosowania procedury (A lub B), gdy bazujemy na wszystkich wyjściowych rankingach.

Tabela 6.

Pozycje obiektów uzyskane w wyniku zastosowania procedur: A, B, wykorzystujące odpowiednio wszystkie wyjściowe rankingi (r_A , r_B) oraz zredukowaną liczbę rankingów (r^*A , r^*B)

Obiekt	Procedura			
	A		B	
	r_A	r^*A	r_B	r^*B
O1	1	1	1	1
O2	7	7	6	7
O3	14	14	13	14
O4	2	2	2	2
O5	6	6	7	6
O6	15	15	15	15
O7	13	13	14	12
O8	10	10	9	10
O9	16	16	16	17
O10	11	11	11	13
O11	4	4	4	4
O12	17	17	17	16
O13	9	9	8	9
O14	8	8	10	8
O15	12	12	12	11
O16	5	5	5	5
O17	3	3	3	3

Źródło: opracowanie własne.

Wyeliminowanie rankingów odstających, a następnie wykorzystanie omówionych procedur (A, B) także nie dało identycznych rankingów, niemniej różnice dotyczyły tylko czterech obiektów.

5. UWAGI KOŃCOWE

Przedstawiona, prawie pół wieku temu przez Hellwiga (1968) metoda porządkowania liniowego obiektów w obszarze ekonomii przyczyniła się niewątpliwie zarówno do rozpowszechniania jak i modyfikowania grupy metod Wielowymiarowej Analizy Porównawczej. Na rozwój tych metod wpływ miała i nadal ma komputeryzacja, ułatwiająca i niejednokrotnie umożliwiająca wykonywanie skomplikowanych obliczeń z dużą dokładnością.

Stojąc przed wyborem metody trzeba zwrócić uwagę na jej własności, charakter zmiennych diagnostycznych, skalę pomiaru zmiennych, a nie stosować aktualnie często wybierane procedury. Tu należy wspomnieć o metodzie pozycyjnej lub TOPSIS, które dały w obu przedstawionych przykładach najbardziej odstające wyniki w porównaniu do rankingów uzyskanych z wykorzystaniem innych metod, co również można zaobserwować w artykułach Kukula, Luty (2015a, 2015b) i Kisieleńska (2016).

Rozważając własności poszczególnych metod normowania cech diagnostycznych warto zwrócić uwagę, iż niektóre z nich w sposób nieliniowy dokonują wartościowania. Jeżeli nie potrafimy sprecyzować wyraźnych przesłanek by realizować przekształcenie cech diagnostycznych w sposób nieliniowy to lepiej poprzestać na liniowym sposobie ich wartościowania.

Z badań przedstawionych w artykule wyciągnąć można następujące wnioski:

- wybierając metodę porządkowania liniowego należy pamiętać, że różne metody dają znacznie różniące się od siebie rankingi,
- dokonując wyboru metody porządkowania liniowego należy przeprowadzić obliczenia z zastosowaniem kilku procedur i dokonać wyboru wykorzystując już istniejące rozwiązania w literaturze lub zweryfikować z ekspertami badanej tematyki,
- zaproponowana procedura eliminacji rankingów „odstających”, a tym samym metod na podstawie, których zostały wyznaczone stanowi, jak się wydaje, pomocne narzędzie przed budową ostatecznego rankingu obiektów.

Ponadto, pragniemy podkreślić, iż odzew w tak krótkim czasie na poruszany przez nas temat w artykule (Kukula, Luty, 2015a), świadczy o tym, że problem ten jest istotny.

LITERATURA

- Bartosiewicz S., (1976), Propozycja metody tworzenia zmiennych syntetycznych, *Zeszyty Naukowe AE*, 84, Wrocław, 5–7.
- Borys T., (1978), Metody normowania cech w statystycznych badaniach porównawczych, *Przegląd Statystyczny*, 25 (2), 227–239.
- Grabiński T., (1984), *Wielowymiarowa analiza porównawcza w badaniach dynamiki zjawisk ekonomicznych*, Zeszyty Naukowe AE w Krakowie, Seria specjalna, Monografie, 61, Kraków.

- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju oraz zasoby i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), 307–327.
- Hwang C. L., Yoon K., (1981), *Multiple Attribute Decision Making: Methods and Applications*, Springer Verlag.
- Kisielińska J., (2016), Ranking państw UE ze względu na potencjalne zaspokojenie zapotrzebowania na produkty rolne z wykorzystaniem metod porządkowania liniowego, *Problemy rolnictwa światowego*, 16 (3), Warszawa, 142–152.
- Kukuła K., (1989), *Statystyczna analiza strukturalna i jej zastosowanie w sferze usług produkcyjnych dla rolnictwa*, Zeszyty Naukowe AE w Krakowie, Seria specjalna, Monografie, 89, Kraków.
- Kukuła K., (2000), *Metoda unitaryzacji zerowanej*, PWN, Warszawa.
- Kukuła K., Luty L., (2015a), Propozycja procedury wspomagającej wybór metody porządkowania liniowego, *Przegląd Statystyczny*, 62 (2), 219–231.
- Kukuła K., Luty L., (2015b), Ranking państw UE ze względu na wybrane wskaźniki charakteryzujące rolnictwo ekologiczne, *Metody ilościowe w Badaniach Ekonomicznych*, Wydawnictwo SGGW, Tom XVI, 3, Warszawa, 225–236.
- Lira J., Wagner W., Wysocki F., (2002), *Mediana w zagadnieniach porządkowania obiektów wielocechowych*, w: Paradyś J., (red.), *Statystyka regionalna w służbie samorządu lokalnego i biznesu*, Internetowa Oficyna Wydawnicza Centrum Statystyki Regionalnej, AE, Poznań, 87–99.
- Młodak A., (2006), *Analiza Taksonomiczna w statystyce regionalnej*, Difin, Warszawa.
- Nowak E., (1977), Syntetyczne mierniki plonów w krajach europejskich, *Wiadomości Statystyczne*, 10, 19–22.
- Pawełek B., (2008), *Metody normalizacji zmiennych w badaniach porównawczych złożonych zjawisk ekonomicznych*, Zeszyty Naukowe UE, Seria specjalna, Monografie, 187, Kraków.
- Panek T., (2009), *Statystyczne metody wielowymiarowej analizy porównawczej*, SGH, Oficyna Wydawnicza, Warszawa.
- Pluta W., (1977), *Wielowymiarowa analiza porównawcza w badaniach ekonomicznych*, PWE, Warszawa.
- Strahl D., (1978), Propozycja konstrukcji miary syntetycznej, *Przegląd Statystyczny*, 25 (2), 205–215.
- Walesiak M., (2014), Przegląd formuł normalizacji wartości zmiennych oraz ich własności w statystycznej analizie wielowymiarowej, *Przegląd Statystyczny*, 61 (4), 363–372.
- Wesołowski W. J., (1975), *Programowanie nowej techniki*, PWN, Warszawa.

JESZCZE O PROCEDURZE WYBORU METODY PORZĄDKOWANIA LINIOWEGO

Streszczenie

W opracowaniu podjęto po raz kolejny zagadnienie wyboru metody porządkowania liniowego jako odpowiedź na propozycję procedury przedstawionej w pracy Kisielińskiej (2016). Pokazano na dwóch przykładach, że ranking zależy od zastosowanej metody, a różnice są istotne. Pokreślono tym samym konieczność zastosowania kilku metod przed budową ostatecznego rankingu oraz wyeliminowaniu rankingów „odstających”. W artykule zaproponowano procedurę ich eliminacji. Oparto ją na średniej arytmetycznej i odchyleniu przeciętnym miar podobieństwa rankingów uzyskanych z wykorzystaniem kilku metod porządkowania liniowego. Z przedstawionych badań wynika, że dokonując wyboru metody porządkowania liniowego należy przeprowadzić obliczenia z zastosowaniem kilku procedur i dokonać wyboru wykorzystując już istniejące rozwiązania w literaturze lub zweryfikować z ekspertami badanej tematyki.

Słowa kluczowe: normowanie cech, wybór metody, porządkowanie liniowe, ranking

ONCE MORE ABOUT THE SELECTION PROCEDURE
FOR THE LINEAR ORDERING METHOD

A b s t r a c t

In this paper once again the issue of selecting the method of linear ordering was analyzed as a response to proposed procedure presented in Kisielińska's paper (2016). It was shown in two examples that the ranking depends on the method used, and the differences are significant. Underlined the need to apply several methods before the construction of the final ranking and eliminating the "outliers" rankings. The article proposes a procedure for their elimination. It is based on the arithmetic mean and deviation of average measures of rankings' similarity obtained with the use of several linear ordering methods. The research shows that when choosing a linear ordering methods calculations should be carried out using several procedures and make a selection using the already existing solutions in the literature or verify with the experts of analyzed subject.

Keywords: characteristics' normalizing, selection of method, linear ordering, ranking

BARTŁOMIEJ LACH¹METODY ŁĄCZENIA I SELEKCJI KLASYFIKATORÓW
W PROGNOZOWANIU UPADŁOŚCI PRZEDSIĘBIORSTW

1. WSTĘP

Ze względu na powszechność i uniwersalność wykorzystania metod klasyfikacyjnych, problem poprawnej klasyfikacji obiektów jest wciąż aktualny i ważny z punktu widzenia badań we wszystkich obszarach nauki. Jedną z dróg dających szansę na poprawę zdolności klasyfikacyjnych obiektów jest jednocześnie wykorzystanie zbioru klasyfikatorów. Podejście to polega na stosowaniu metod, w których ostateczna decyzja o przynależności obiektu do jednej z możliwych populacji (klas) zależy od wskazań większej liczby klasyfikatorów indywidualnych. W niniejszym artykule, klasyfikatorami indywidualnymi nazywane są metody klasyfikacyjne różnego typu, których prognozy (wskazania klasy przynależności lub prawdopodobieństwa a posteriori) służą do podejmowania ostatecznej decyzji o przynależności obiektu przez cały zespół klasyfikatorów. Zbiór klasyfikatorów podejmujący decyzje klasyfikacyjne zgodnie z regułami metod łączenia lub selekcji nazywany jest w skrócie klasyfikatorem hybrydowym. Modelami klasyfikacyjnymi indywidualnymi lub hybrydowymi nazywane są natomiast konkretne realizacje metod klasyfikacyjnych, oszacowane na podstawie obiektów próby uczącej i testowane na obiektach próby testującej.

W niniejszym artykule przedstawione zostały wyniki badań nad próbą wykorzystania metod łączenia oraz selekcji klasyfikatorów w prognozowaniu upadłości przedsiębiorstw. Poza porównaniem trafności prognoz uzyskanych przez modele bazujące na klasyfikatorach indywidualny oraz hybrydowych, przeanalizowano także wpływ liczby uwzględnianych zmiennych objaśniających na poprawność klasyfikacji przez poszczególne modele. W zbiorze klasyfikatorów indywidualnych znalazły się cztery różne metody klasyfikacyjne: liniowa analiza dyskryminacyjna, regresja logistyczna, sztuczna sieć neuronowa oraz las losowy.

Pojęcie „upadłości przedsiębiorstwa” jest silnie powiązane z pojęciem „bankructwa przedsiębiorstwa”. Pomimo częstego zamiennego używania tych pojęć w języku potocznym, należy podkreślić, że nie są to określenia tożsame. Termin bankructwa podkreśla ekonomiczny charakter zjawiska, jakim jest zaprzestanie podstawowej działalności przedsiębiorstwa. Upadłość jest natomiast pojęciem prawnym, którego znacze-

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Ekonometrii, al. Niepodległości 10, 61-875 Poznań, Polska, e-mail: lach.bartlomiej@gmail.com.

nie regulują przepisy prawa upadłościowego. W niniejszym artykule, autor korzysta z pojęcia „upadłości przedsiębiorstw” ze względu na wykorzystanie w konstruowanych modelach klasyfikacyjnych zmiennej kategoryjnej określającej fakt wystąpienia lub nie wystąpienia upadłości w sensie prawnym.

Analiza dyskryminacyjna z liniową funkcją dyskryminacyjną została zaproponowana przez Fishera (1936) jako narzędzie pozwalające odseparować w możliwie najlepszy sposób obiekty pochodzące z dwóch różnych populacji. Pierwszym modelem dyskryminacyjnym wykorzystanym do problemu prognozowania upadłości przedsiębiorstw był model Altmanna nazwany „*Z-score model*” (Altman, 1968). Profesor Altman jest także autorytetem w dziedzinie teorii bankructwa oraz autorem wielu publikacji na ten temat (Altman i inni, 2008; Altman, Roggi, 2013). Modele liniowej analizy dyskryminacyjnej są narzędziem często wykorzystywanym w prognozowaniu upadłości przedsiębiorstw zarówno zagranicą jak i w Polsce (Mossman i inni, 1998; Mączyńska, 1994; Hadasik, 1998; Hołda, 2001, 2006; Mączyńska, Zawadzki, 2006). Proponowane modele wykorzystują różne zestawy wskaźników lub wielkości finansowych, które pozwalają z dużym prawdopodobieństwem wskazywać przedsiębiorstwa zagrożone ryzykiem upadłości.

Drugim z wykorzystanych klasyfikatorów indywidualnych jest regresja logistyczna. Metoda ta pozwala określać przynależność obiektu do jednej z dwóch populacji na podstawie zmiennych o charakterze ilościowym bądź jakościowym. Przynależność wynika z wielkości prawdopodobieństwa modelowanego z wykorzystaniem funkcji przekształcającej prawdopodobieństwo na logarytm ilorazu szans (logit) oraz dystrybuanty rozkładu logistycznego (Pociecha i inni, 2014). Warto podkreślić, że modele regresji logistycznej nie wymagają spełnienia założenia o normalności rozkładu zmiennych objaśniających. Metoda ta, podobnie jak liniowa analiza dyskryminacyjna, jest często wykorzystywana w prognozowaniu upadłości przedsiębiorstw (Ohlson, 1980; Hołda, 2000; Wędzki, 2005).

Kolejnym z wykorzystanych w badaniu klasyfikatorów indywidualnych jest sztuczna sieć neuronowa, której działanie wzorowane jest na sieci neuronów przetwarzających sygnały w układzie nerwowym organizmów żywych. Sztuczna sieć neuronowa poprzez sieć powiązanych neuronów stanowiących określoną strukturę ma za zadanie przetworzyć wejściowe sygnały w taki sposób, aby sygnały wychodzące z sieci były zgodne z oczekiwanymi. Obok wyboru odpowiedniej struktury, ważnym etapem konstruowania sieci jest proces jej uczenia z wykorzystaniem odpowiednich algorytmów. W badaniach nad upadłością przedsiębiorstw najczęściej stosowaną postacią sieci neuronowej jest perceptron wielowarstwowy (Bell i inni, 1990; Korol, Prusak, 2005; Hołda, 2006).

Ostatnim klasyfikatorem indywidualnym uwzględnionym przez autora artykułu jest las losowy. Podobnie jak sztuczna sieć neuronowa, las losowy należy do grupy metod nieparametrycznych. Nie wymaga zatem spełnienia rygorystycznych założeń wobec zmiennych objaśniających wykorzystywanych przy jego budowie. Las zbudowany jest z dużej liczby drzew klasyfikacyjnych, które poprzez głosowanie większościowe

wskazują na populację (klasę), do której należy zakwalifikować rozpoznawany obiekt. Należy zauważyć, że decyzje klasyfikacyjne lasu losowego bazują na wskazaniach większej liczby drzew klasyfikacyjnych. Celem autora jest zweryfikowanie użyteczności metod łączenia i selekcji dla metod różnego typu, dlatego las losowy został włączony do grupy klasyfikatorów indywidualnych. Skuteczność lasu losowego w prognozowaniu upadłości przedsiębiorstw była weryfikowana również dla przedsiębiorstw działających w Polsce (Korol, 2010).

2. METODY ŁĄCZENIA I SELEKCJI KLASYFIKATORÓW

Ze względu na odmienne koncepcje działania metod klasyfikacyjnych uwzględnionych w badaniu, może się zdarzyć, że prognozy uzyskiwane na podstawie modeli uczonych z wykorzystaniem tego samego zbioru obiektów istotnie różnią się od siebie w zależności od wykorzystanej metody. W sytuacji kiedy rozpatrywane metody dokonują błędnej klasyfikacji odmiennych obiektów, a żadnej z metod nie można jednoznacznie wskazać jako najlepszej, zasadna jest próba jednoczesnego wykorzystania metod indywidualnych stosując podejście łączenia bądź selekcji. Klasyfikatory indywidualne powinny zatem charakteryzować się wysoką precyzją oraz możliwą różnorodnością stawianych prognoz² (Kuncheva, 2000). Stosując metody łączenia lub selekcji, ostateczna decyzja o przynależności obiektu do populacji podejmowana jest z wykorzystaniem większej liczby metod klasyfikacyjnych, dając tym samym szansę na uzyskanie pełniejszej informacji o klasyfikowanych obiektach oraz poprawę wyników ogólnej trafności (Krzyśko i inni, 2008).

Łączenie klasyfikatorów jest sposobem wykorzystania wielu metod klasyfikacyjnych jednocześnie. Ostateczna decyzja o przynależności obiektu do jednej ze zdefiniowanych klas podejmowana jest na podstawie wskazań klasyfikatorów indywidualnych (metoda głosowania większościowego) lub na podstawie wektorów prawdopodobieństw a posteriori przynależności obiektów (np. metoda sumacyjna).

W metodzie głosowania większościowego decyzja o przynależności obiektu zależy od liczby „głosów” oddanych na każdą z możliwych klas. Zaletą głosowania większościowego jest możliwość zastosowania dla dowolnych metod klasyfikacyjnych wskazujących jedynie klasę, do której zaliczany jest obiekt. Nie jest wymagana znajomość wielkości prawdopodobieństwa a posteriori przynależności obiektu do danej klasy. Należy mieć na uwadze, że w przypadku parzystej liczby klasyfikatorów składowych oraz dwóch klas, z których mogą pochodzi obiekty, metoda może nie rozstrzygać o przynależności obiektu. Dzieje się tak w sytuacji gdy liczba głosów oddanych na każdą z klas jest taka sama (Krzyśko i inni, 2008).

Zakładając, że występuje c różnych klasyfikatorów indywidualnych $\hat{d}_1, \hat{d}_2, \dots, \hat{d}_c$, wyznaczonych na bazie n -elementowej próby uczącej $U_n = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)\}$,

² Wysoka różnorodność jest rozumiana jako możliwie niska lub nawet ujemna korelacja błędów poszczególnych klasyfikatorów indywidualnych.

gdzie $\mathbf{x}_i$ jest wektorem obserwowanych cech, a y_1 etykietą jednej z K klas, klasyfikator głosowania można zapisać jako³:

$$\hat{d}_V(\mathbf{x}) = \arg \max_{1 \leq k \leq K} \sum_{j=1}^c I(\hat{d}_j(\mathbf{x}) = k), \quad j = 1, 2, \dots, c. \quad (1)$$

Innym sposobem łączenia klasyfikatorów jest metoda sumacyjna. Warunkiem możliwości wykorzystania metody sumacyjnej jest posiadanie ocen prawdopodobieństw a posteriori przynależności klasyfikowanego obiektu do zdefiniowanych klas $\hat{p}^{(j)}(k|\mathbf{x})$, $j = 1, 2, \dots, c$; $k = 1, 2, \dots, K$. Klasyfikator sumacyjny można wtedy zapisać:

$$\hat{d}_{SUM}(\mathbf{x}) = \arg \max_k s_k(\mathbf{x}), \quad (2)$$

gdzie:

$$s_k(\mathbf{x}) = \sum_{j=1}^c \hat{p}^{(j)}(k|\mathbf{x}).$$

Sumę prawdopodobieństw można zastąpić również ich wartością średnią uzyskując takie same wskazania klasyfikatora. Możliwe jest także wprowadzenie dodatkowych wag różnicujących wpływ określonych klasyfikatorów na podejmowaną ostatecznie decyzję. Przykładem mogą być wagi proporcjonalne do współczynników trafności globalnej klasyfikatorów indywidualnych.

Metoda iloczynowa, w przeciwieństwie do metody sumacyjnej, wykorzystuje iloczyny prawdopodobieństw a posteriori. Z tego powodu, w przypadku gdy chociaż jeden z klasyfikatorów indywidualnych wskazuje niskie prawdopodobieństwo przynależności obiektu do danej klasy to znacząco obniża to prawdopodobieństwo tego, że klasyfikator łączony wskaże ostatecznie tę klasę dla rozważanego obiektu. Klasyfikator iloczynowy można zapisać jako:

$$\hat{d}_{PROD}(\mathbf{x}) = \arg \max_k r_k(\mathbf{x}), \quad (3)$$

gdzie:

$$r_k(\mathbf{x}) = \prod_{j=1}^c \hat{p}^{(j)}(k|\mathbf{x}).$$

Kolejnym sposobem łączenia klasyfikatorów jest wykorzystanie metod bazujących na statystykach pozycyjnych wyznaczanych dla prawdopodobieństw a posteriori przynależności obiektów do poszczególnych klas. Łączenie klasyfikatorów odbywa się dla statystyk: minimum, maksimum oraz mediana, co można zapisać:

³ $I(q)$ jest funkcją wskaźnikową przyjmującą wartość 1 gdy zdanie q jest prawdziwe lub 0 gdy zdanie q jest fałszywe.

$$\text{Metoda minimum } \hat{d}_{MIN}(\mathbf{x}) = \arg \max_k \min_j \hat{p}^{(j)}(k|\mathbf{x}), \quad (4)$$

$$\text{Metoda maksimum } \hat{d}_{MAX}(\mathbf{x}) = \arg \max_k \max_j \hat{p}^{(j)}(k|\mathbf{x}), \quad (5)$$

$$\text{Metoda medianowa } \hat{d}_{ME}(\mathbf{x}) = \arg \max_k Me \hat{p}^{(j)}(k|\mathbf{x}). \quad (6)$$

Drugą z grup metod wykorzystujących zbiór klasyfikatorów indywidualnych do określania przynależności obiektu jest selekcja. Selekcja klasyfikatorów zakłada, że poszczególne metody klasyfikacji mogą działać skuteczniej, kiedy stosuje się je do obiektów podobnych ze względu na pewne cechy. Możliwy jest zatem wybór lokalnie najlepszych metod klasyfikacyjnych, dla których łączna trafność jest wyższa niż każdej metody z osobna (Dasarathy, Sheela, 1979).

Jedną z propozycji wykorzystania metody selekcji przedstawiła Kuncheva (2000). Badaczka zaproponowała wykorzystanie metody iteracyjnej optymalizacji podziału zbioru obiektów – k-średnich w celu wyznaczenia skupień obiektów, z którymi następnie skojarzone zostały najskuteczniejsze w tych skupieniach klasyfikatory indywidualne. Autorka wykorzystała w badaniu pięć oddzielnie nauczonych sieci neuronowych. Inne metody łączenia klasyfikatorów zostały dokładnie opisane w pracy Kunchevej (2004).

3. CELE BADANIA, PROCEDURA POSTĘPOWANIA ORAZ WYKORZYSTANY MATERIAŁ BADAWCZY

Przedstawione powyżej koncepcje łączenia oraz selekcji klasyfikatorów posłużyły do budowy modeli prognozowania upadłości przedsiębiorstw bazujących na wskazaniach czterech różnych klasyfikatorów indywidualnych: liniowa analiza dyskryminacyjna, regresja logistyczna, sztuczna sieć neuronowa oraz las losowy. Powstałe w ten sposób modele hybrydowe porównywane były z modelami indywidualnymi pod względem trafności stawianych przez nie prognoz.

Podstawowym celem badania było wskazanie metod klasyfikacyjnych spośród wszystkich analizowanych (4 klasyfikatory indywidualne + 9 klasyfikatorów łączenia lub selekcji) zapewniających możliwie wysoką jakość prognozowania upadłości przedsiębiorstw dla obiektów próby testującej. Badanie służy weryfikacji hipotezy mówiącej o wyższej skuteczności poprawnej klasyfikacji obiektów przez modele bazujące na metodach łączenia lub selekcji klasyfikatorów w porównaniu z modelami indywidualnymi. Istotnym elementem badania było także wskazanie listy wskaźników finansowych, których wykorzystanie w konstrukcji modeli klasyfikacyjnych zapewniło najwyższą jakość stawianych prognoz. Badanie umożliwiło również przeprowadzenie analizy skuteczności metod klasyfikacji w zależności od liczby wskaźników finansowych wykorzystywanych do budowy poszczególnych modeli indywidualnych oraz hybrydowych. Wyniki podjętej analizy mogą pomóc w odpowiedzi na pytanie, czy i kiedy warto stosować metody łączenia oraz selekcji klasyfikatorów do problemu prognozowania upadłości przedsiębiorstw.

Tabela 1.

Klasyfikatory indywidualne i hybrydowe wykorzystane w badaniu

Oznaczenie	Klasyfikator
AD	Liniowa analiza dyskryminacyjna
RL	Regresja logistyczna
SSN	Sztuczna sieć neuronowa
LL	Las losowy
H1	Głosowanie większościowe
H2	Klasyfikator sumacyjny
H3	Klasyfikator sumacyjny (ważony)
H4	Klasyfikator iloczynowy
H5	Klasyfikator minimum
H6	Klasyfikator medianowy
H7	Klasyfikator maksimum
H8	Klasyfikator selekcji 1
H9	Klasyfikator selekcji 2

Źródło: opracowanie własne.

Konstruowane w badaniu sztuczne sieci neuronowe miały postać perceptronu wielowarstwowego. Sieci składały się za każdym razem z 3 warstw (1 wejściowa, 1 ukryta oraz 1 wyjściowa). Liczba neuronów warstwy wejściowej odpowiadała liczbie wskaźników finansowych wykorzystanych do budowy poszczególnych sieci. Liczba neuronów warstwy ukrytej stanowiła za każdym razem połowę liczby neuronów warstwy wejściowej. W przypadku nieparzystej liczby neuronów pierwszej warstwy, liczba ta była zaokrąglana w górę. Reguły związane z ustaleniem liczby neuronów w każdej z warstw sieci związane są z automatyzacją procesu budowy dużej liczby SSN w dalszej części badania. Uczenie sieci odbywało się z wykorzystaniem metody zmiennej metryki – algorytm BFGS (Krzyśko i inni, 2008).

W przeprowadzonym badaniu zbudowano także dużą liczbę lasów losowych. Do budowy pojedynczego lasu wykorzystywano za każdym razem 1000 drzew klasyfikacyjnych. W każdym węźle drzewa, do podziału obiektów wybierana była jedna z co najwyżej czterech losowych zmiennych uczestniczących w budowie lasu. Miarą różnorodności klas w węźle przy budowie drzew klasyfikacyjnych był indeks Giniego.

Dodatkowego komentarza wymaga także wykorzystanie niektórych klasyfikatorów hybrydowych. W przypadku klasyfikatora sumacyjnego (ważonego) wykorzystano wagi proporcjonalne do współczynników trafności globalnej uzyskanych na poziomie próby uczącej dla klasyfikatorów indywidualnych.

W przypadku klasyfikatora selekcji 1, lokalnie najlepszy klasyfikator wyznaczany był dla każdego z 4 skupień obiektów podobnych pod względem opisujących je cech diagnostycznych (zestaw wskaźników finansowych). Do utworzenia 4 skupień

wykorzystano metodę k-średnich. W przypadku klasyfikatora selekcji 2, podobieństwo obiektów pod względem opisujących je cech diagnostycznych zastąpiono podobieństwem pod względem wielkości prawdopodobieństw a posteriori przynależności do populacji spółek upadłych (prawdopodobieństwa uzyskane na podstawie 4 klasyfikatorów indywidualnych).

W badaniu empirycznym wykorzystano informacje dotyczące sytuacji finansowo-majątkowej 180 spółek akcyjnych z okresu 1999–2012 działających w Polsce w sektorach: budownictwo, handel oraz przetwórstwo przemysłowe. Dla każdego z sektorów dysponowano 30 obiektami, wobec których odpowiedni sąd ogłosił upadłość podmiotu gospodarczego oraz 30 obiektami w dobrej kondycji finansowej, których wielkość aktywów oraz typ prowadzonej działalności odpowiadały poszczególnym spółkom, które ogłosiły upadłość. Poza zmienną zero-jedynkową wskazującą na wystąpienie bądź nie występowanie upadłości przedsiębiorstwa, obiekty badania opisane zostały zestawem 19 wskaźników finansowych prezentujących w możliwie pełny sposób ich sytuację finansowo-majątkową. W przypadku spółek upadłych, wskaźniki finansowe wyznaczane były na podstawie sprawozdań finansowych z okresu na rok przed złożeniem pierwszego wniosku o ogłoszenie upadłości. Źródłami informacji i danych w badaniu były: Internetowy Monitor Sądowy i Gospodarczy, Monitor Polski B oraz baza danych firmy Notoria Serwis.

Tabela 2.

Lista wskaźników finansowych wykorzystanych w badaniu

Zmienna	Wskaźnik finansowy
X1	Stopa zwrotu z aktywów = $\text{zysk netto} / \text{średnia wartość aktywów}$
X2	Stopa zwrotu z kapitału własnego = $\text{zysk netto} / \text{średnia wartość kapitałów własnych}$
X3	Zysk brutto / średnia wartość aktywów
X4	Zysk ze sprzedaży / średnia wartość aktywów
X5	Marża zysku brutto = $\text{zysk brutto} / \text{przychody ze sprzedaży}$
X6	Marża zysku netto = $\text{zysk netto} / \text{przychody ze sprzedaży}$
X7	Marża zysku operacyjnego = $\text{zysk na działalności operacyjnej} / \text{przychody ze sprzedaży}$
X8	$(\text{Aktywa obrotowe} - \text{zobowiązania krótkoterminowe}) / \text{suma bilansowa}$
X9	Wskaźnik bieżącej płynności = $\text{Aktywa obrotowe} / \text{zobowiązania krótkoterminowe}$
X10	Wskaźnik szybkiej płynności = $(\text{Aktywa obrotowe} - \text{zapasy}) / \text{zobowiązania krótkoterminowe}$
X11	Wskaźnik podwyższonej płynności = $(\text{Aktywa obrotowe} - \text{zapasy} - \text{należności}) / \text{zobowiązania krótkoterminowe}$
X12	$\text{Zobowiązania ogółem} / \text{aktywa ogółem}$
X13	$\text{Zobowiązania długoterminowe} / \text{aktywa ogółem}$
X14	$\text{Kapitał własny} / \text{aktywa ogółem}$
X15	$\text{Kapitał własny} / \text{zobowiązania ogółem}$

Tabela 2. (cd.)

Zmienna	Wskaźnik finansowy
X16	Wskaźnik rotacji należności = średnia wartość należności / przychody ze sprzedaży netto · 365
X17	Wskaźnik rotacji zapasów = średnia wartość zapasów / przychody ze sprzedaży netto · 365
X18	Wskaźnik rotacji zobowiązań = średnia wartość zobowiązań / przychody ze sprzedaży · 365
X19	Wskaźnik rotacji aktywów = średnia wartość aktywów / przychody ze sprzedaży · 365

Źródło: opracowanie własne.

Nadrzędnym celem badania była próba generalnej oceny przydatności metod łączenia i selekcji klasyfikatorów w prognozowaniu upadłości przedsiębiorstw oraz chęć przedstawienia wpływu różnej liczby zmiennych objaśniających na skuteczność działania poszczególnych metod. Z tego powodu, w toku prowadzonych badań, konstruowana była duża liczba modeli klasyfikacyjnych z różnym co do liczebności oraz składu zestawem zmiennych objaśniających. W pierwszym kroku badania, dla każdej z możliwych liczebności zmiennych w modelach (przyjęto od 2 do 18) wylosowano po 100 różnych kombinacji zmiennych spośród wszystkich 19 dostępnych wskaźników finansowych⁴. Autor badania ma świadomość, że w losowym zestawie zmiennych mogą znaleźć się pary zmiennych silnie skorelowanych. Taka analiza również została przeprowadzona. Nie zdecydowano się jednak na usuwanie zmiennych skorelowanych, ponieważ nadrzędnym kryterium oceny przyjętym w badaniu była zdolność poprawnej klasyfikacji obiektów próby testującej, a nie np. własności statystyczno-ekonometryczne konstruowanych modeli (szczególnie parametrycznych). Rozważania na ten temat w szerszym zakresie podjął Zbigniew Czerwiński w książce pt. „Moje zmagania z ekonomią” (Czerwiński, 2002).

Dla każdej wylosowanej kombinacji zmiennych szacowane były modele indywidualne (AD, RL, SSN, LL). Następnie, na podstawie wskazań oraz wektorów prawdopodobieństw a posteriori przynależności obiektów do populacji pochodzących z modeli indywidualnych zbudowano modele hybrydowe (H1, H2, H3, H4, H5, H6, H7, H8, H9). Za każdym razem modele były szacowane i walidowane poprzez 25-krotne⁵ wylosowanie próby testującej (losowanie proste bez zwracania) stanowiącej 30% wszystkich dostępnych obserwacji. Pozostałe 70% obserwacji stanowiło próbę uczącą. Dla każdego pojedynczego podziału zbioru obiektów na próbę uczącą i testującą wyznaczano wszystkie 13 modeli klasyfikacyjnych. Uzyskane wyniki współczynników trafności globalnej prognoz w próbie uczącej i testującej były uśredniane dla 25 powtórzeń. W procesie walidacji, dla każdej pojedynczej wylosowanej kombinacji zmiennych

⁴ Wyjątkami były losowania 2-elementowych oraz 18-elementowych kombinacji spośród 19 wskaźników finansowych. W tych przypadkach zbadano wszystkie możliwe kombinacje dokonując przeglądu pełnego możliwych zestawów zmiennych w modelach klasyfikacyjnych. Było ich odpowiednio 171 oraz 19.

⁵ Ustalona liczba powtórzeń równa 25 wynika z przeprowadzonej dodatkowej analizy stabilności współczynników trafności globalnej prognoz dla analizowanych metod.

wyznaczono zatem 325 modeli indywidualnych i hybrydowych $((4 + 9) \cdot 25)$, a łączna liczba modeli uwzględniona w analizie wyniosła 549 250.

Wszystkie analizy wykonano z wykorzystaniem autorskiej aplikacji komputerowej napisanej w środowisku R, mającej zastosowanie dla problemów binarnej klasyfikacji obiektów.

4. BADANIE SKUTECZNOŚCI WYKORZYSTANIA METOD ŁĄCZENIA ORAZ SELEKCJI KLASYFIKATORÓW W PROGNOZOWANIU UPADŁOŚCI PRZEDSIĘBIORSTW

W tabeli 3 przedstawiono zestawienie modeli klasyfikacyjnych z różną liczbą zmiennych objaśniających, których średni współczynnik trafności globalnej⁶ wyznaczony dla obiektów próby testującej był najwyższy.

Tabela 3.

Modele z najwyższymi współczynnikami trafności globalnej w próbie testującej

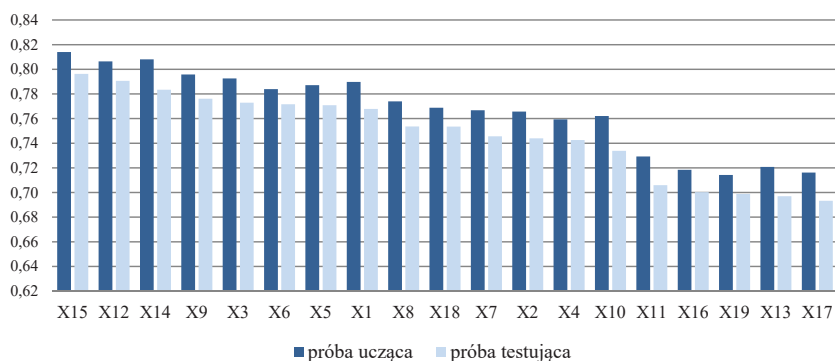
Liczba zmiennych	Wsp. trafności globalnej	Metoda	Zestaw zmiennych
2	0,847	RL	X5, X12
3	0,848	RL	X1, X2, X15
4	0,865	RL	X5, X10, X11, X12
5	0,869	LL	X2, X4, X10, X14, X16
6	0,869	H2	X5, X6, X12, X14, X17, X18
7	0,852	LL	X2, X6, X9, X14, X16, X18, X19
8	0,857	H2 i H6	X1, X2, X5, X7, X8, X15, X17, X18
9	0,859	H2	X3, X4, X5, X8, X9, X10, X12, X18, X19
10	0,859	H6	X4, X7, X9, X10, X11, X13, X14, X16, X17, X18
11	0,858	LL	X2, X3, X5, X8, X9, X10, X13, X14, X16, X18, X19
12	0,866	LL	X2, X3, X4, X5, X6, X7, X8, X11, X12, X13, X17, X18
13	0,858	LL	X1, X3, X4, X6, X7, X8, X9, X10, X14, X15, X16, X18, X19
14	0,869	LL	X2, X3, X4, X5, X8, X9, X11, X13, X14, X15, X16, X17, X18, X19
15	0,863	LL	X4, X5, X6, X7, X8, X9, X10, X11, X12, X13, X14, X16, X17, X18, X19
16	0,862	LL	X2, X3, X4, X5, X6, X8, X9, X10, X12, X13, X14, X15, X16, X17, X18, X19
17	0,867	LL	X2, X3, X4, X6, X7, X8, X9, X10, X11, X12, X13, X14, X15, X16, X17, X18, X19
18	0,853	LL	X1, X2, X3, X4, X5, X6, X7, X8, X9, X10, X11, X13, X14, X15, X16, X17, X18, X19

Źródło: opracowanie własne.

⁶ W wyniku przeprowadzonej walidacji, wartości współczynnika trafności globalnej zostały uśrednione dla 25 powtórzeń.

Na podstawie informacji zawartych w tabeli 3 można stwierdzić, że w przypadku małej liczby zmiennych (od 2 do 4) najwyższe zdolności poprawnej klasyfikacji obiektów uzyskały modele regresji logistycznej. W przypadku modeli z dużą liczbą zmiennych (od 11 do 18) za każdym razem najskuteczniejszy okazał się las losowy. Modele oparte na metodach łączenia H2 (klasyfikator sumacyjny) oraz H6 (klasyfikator medianowy) uzyskiwały najwyższe wyniki dla liczby zmiennych 6 oraz od 8 do 10.

Zbudowanie modeli dla wszystkich możliwych 2-elementowych kombinacji zmiennych objaśniających umożliwiło zbadanie w sposób numeryczny zdolności poprawnego klasyfikowania obiektów przez poszczególne wskaźniki finansowe wykorzystane w badaniu. Na rysunku 1 przedstawiono średnie współczynniki trafności globalnej w próbie uczącej oraz testującej dla wszystkich modeli z dwoma zmiennymi objaśniającymi, w których wystąpiła każda z analizowanych zmiennych (każdy ze wskaźników finansowych wystąpił w 18 kombinacjach zestawów zmiennych).



Rysunek 1. Zdolności poprawnej klasyfikacji obiektów przez wskaźniki finansowe

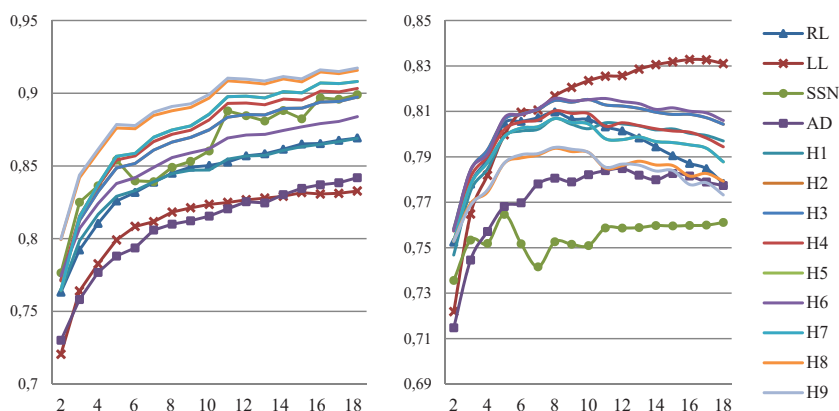
Źródło: opracowanie własne.

Najwyższe współczynniki trafności globalnej uzyskiwały modele, w których znajdowały się zmienne dotyczące zadłużenia przedsiębiorstwa (X15, X12), relacji kapitału własnego do wielkości aktywów (X14) oraz bieżącej płynności (X9). Najniższe współczynniki odpowiadały modelom, w których wystąpiły wskaźniki rotacji należności, zapasów i aktywów (X16, X17, X19).

Na rysunku 2 przedstawiono kształtowanie się średnich współczynników trafności globalnej uzyskanych przez poszczególne modele w próbie uczącej (lewy wykres) oraz testującej (prawy wykres). Dodatkowo na wykresie wyróżniono wyniki uzyskane przez modele indywidualne. Warto przypomnieć, że dla każdej z możliwych liczebności zmiennych (od 2 do 18) wylosowano po 100 kombinacji zmiennych⁷, dla których 25-krotnie wyznaczano i testowano wszystkie 13 modeli klasyfikacyjnych

⁷ Wyjątek stanowiły modele dla 2 i 18 zmiennych.

(indywidualne i hybrydowe). Uzyskane rezultaty pokazują, że w przypadku próby uczącej, średnia wartość współczynnika trafności globalnej dla wszystkich analizowanych klasyfikatorów wzrastała wraz ze wzrostem liczby zmiennych. Najskuteczniejsze w próbie uczącej okazały się klasyfikatory selekcji uzyskujące średnią trafność globalną przekraczającą 90%. Najlepszym z klasyfikatorów indywidualnych (dla różnych liczebności zmiennych) była sztuczna sieć neuronowa.



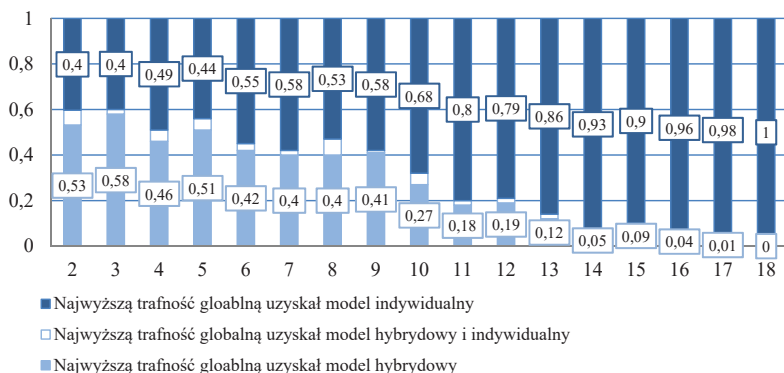
Rysunek 2. Liczba zmiennych w modelach a trafność prognoz w próbach uczącej i testującej

Źródło: opracowanie własne.

Dla zdecydowanej większości klasyfikatorów, budowane modele zwiększały zdolności poprawnej klasyfikacji obiektów próby testującej do momentu uwzględniania liczby zmiennych nieprzekraczającej 8–10. Przy wyższej liczbie zmiennych współczynniki trafności globalnej stabilizowały się lub spadały. Wyjątkiem jest klasyfikator lasu losowego, dla którego średnia trafność prognoz w próbie testującej wzrastała razem ze wzrostem liczby uwzględnianych zmiennych. Las losowy uzyskiwał średnio najwyższą skuteczność prognoz w modelach z liczbą zmiennych większą od 8. Przy mniejszej liczbie zmiennych, najwyższe wyniki uzyskiwały modele dla klasyfikatorów hybrydowych (sumacyjny i medianowy) oraz modele regresji logistycznej. Należy zauważyć duży spadek trafności prognoz w przypadku sztucznych sieci neuronowych. W celu poprawy uzyskiwanych wyników możliwa jest potrzeba zmiany struktury konstruowanych sieci lub przyjętego algorytmu uczenia. Takie same wyniki uzyskały pary modeli dla klasyfikatorów hybrydowych: sumacyjny i sumacyjny (ważony), a także minimum i maksimum.

Na rysunku 3 przedstawiono dla każdej z możliwych liczebności zmiennych odsetek modeli indywidualnych oraz hybrydowych, które uzyskały najwyższą trafność globalną prognoz w próbie testującej. Dla przykładu, pośród 100 kombinacji modeli z 4 zmiennymi objaśniającymi, w 49 przypadkach modelem, który uzyskał najwyższą trafność globalną w próbie testującej (po walidacji) był jeden z 4 modeli indywidualnych. W 46 przypadkach natomiast, najwyższą trafność globalną wskazał jeden

z 9 modeli hybrydowych. W pozostałych 5 przypadkach najlepszy z modeli indywidualnych oraz hybrydowych wskazały taką samą wielkość współczynnika trafności globalnej.



Rysunek 3. Odsetek modeli z najwyższym współczynnikiem trafności globalnej w próbie testującej a liczba uwzględnianych zmiennych

Źródło: opracowanie własne.

Pośród modeli z małą liczbą zmiennych (2, 3, 5), w ponad połowie przypadków przynajmniej jeden z modeli hybrydowych uzyskiwał wyższy współczynnik trafności globalnej na poziomie próby testującej niż najlepszy z modeli indywidualnych. Odsetek ten spada wraz ze wzrostem liczby zmiennych. Sytuacja ta jest związana z rosnącą (wraz ze wzrostem liczby zmiennych) skutecznością poprawnej klasyfikacji obiektów przez las losowy należący do grupy klasyfikatorów indywidualnych. Wzrost ten można zaobserwować na rysunku 2. Warto przypomnieć, że niektóre modele hybrydowe (H2 i H6) znalazły się również w zestawieniu najlepszych modeli zaprezentowanych w tabeli 3.

5. PODSUMOWANIE

Przeprowadzone badanie empiryczne pozwoliło na realizację założonego celu jakim była ocena skuteczności wykorzystania metod łączenia oraz selekcji klasyfikatorów w prognozowaniu upadłości przedsiębiorstw. Badanie zostało zaprojektowane z myślą o możliwie wysokiej porównywalności uzyskiwanych wyników dla różnych metod klasyfikacyjnych. Dokonana analiza pozwala stwierdzić, że metody łączenia oraz selekcji, bazujące na wskazaniach oraz prawdopodobieństwach a posteriori klasyfikatorów indywidualnych, mogą wpływać na poprawę jakości stawianych prognoz. W badaniu wykorzystano 4 klasyfikatory indywidualne oraz 9 klasyfikatorów hybrydowych. W tabeli 3 przedstawiającej zestawienie modeli o najwyższych zdolnościach

poprawnego klasyfikowania obiektów kilkakrotnie znalazły się modele hybrydowe (sumacyjny oraz medianowy).

W toku realizacji podstawowego celu badania wskazano również zestaw modeli klasyfikacyjnych dla różnej liczby zmiennych objaśniających, które charakteryzowały się najwyższymi wynikami poprawnej klasyfikacji obiektów próby testującej (tabela 3). Oceniając uzyskane wyniki klasyfikacji warto podkreślić, że spółki uwzględnione w badaniu pochodziły z trzech różnych branż, co mogło być istotnym czynnikiem utrudniającym poprawną klasyfikację obiektów charakteryzujących się różnorodną specyfiką prowadzonej działalności. Na podstawie przeglądu wyników klasyfikacji uzyskanych przez innych autorów badań wynika, że średnia zdolność poprawnej klasyfikacji (w próbie testującej) modeli szacowanych dla przedsiębiorstw działających w Polsce w branży produkcyjnej mieści się w przedziale 77–85%, w branży budowlanej: 65–90%, a w branży handlowo-usługowej: 61–82% (Pociecha i inni, 2014). Wyniki poprawnej klasyfikacji najlepszych z uzyskanych modeli w badaniu, sięgające 87% (w próbie testującej) można zatem uznać za względnie wysokie.

W toku prowadzonej analizy dokonano również oceny dostępnych wskaźników finansowych pod względem zdolności poprawnego klasyfikowania obiektów do populacji spółek upadłych oraz tych w dobrej kondycji finansowej. Najwyższe oceny uzyskały wskaźniki zadłużenia przedsiębiorstwa, bieżącej płynności oraz relacji kapitału własnego do wielkości aktywów. Najniżej ocenione zostały natomiast wskaźniki rotacji należności, zapasów oraz aktywów.

Analiza prognoz stawianych przez dużą liczbę modeli klasyfikacyjnych, uwzględniających różną liczbę zmiennych objaśniających pozwala stwierdzić, że stosowanie metod łączenia oraz selekcji klasyfikatorów jest częściej skuteczne w przypadku modeli z mniejszą liczbą zmiennych. Wniosek ten wynika z porównania rezultatów klasyfikacji obiektów i dotyczy zbioru metod uwzględnionych w badaniu. Pośród wszystkich rozważanych metod klasyfikacyjnych, najwyższą jakością stawianych prognoz, w sytuacji uwzględnienia dużej liczby zmiennych objaśniających, charakteryzował się las losowy.

LITERATURA

- Altman E. I., (1968), Financial Ratios, Discriminant Analysis and Prediction of Corporate Bankruptcy, *The Journal of Finance*, 23 (4), 589–609.
- Altman E. I., Roggi O., (2013), *Managing and Measuring Risk. Emerging Global Standards and Regulations After Financial Crisis*, World Scientific Press.
- Altman E. I., Caouette J. B., Narayanan P., Nimmo R., (2008), *Managing Credit Risk: The Great Challenge for Global Financial Markets*, Wiley Finance.
- Bell T. B., Ribar G. S., Verchio J., (1990), *Neural Nets Versus Logistic Regression: A comparison of Each Model's Ability to Predict Commercial Bank Failures*, Proceedings of the 1990 Deloitte and Touche/University of Kansas Symposium on Auditing Problems, 29–53.
- Czerwiński Z., (2002), *Moje zmagania z ekonomią*, Wydawnictwo Akademii Ekonomicznej w Poznaniu, Poznań.

- Dasarathy B. V., Sheela B. V., (1979), *A Composite Classifier System Design: Concepts and Methodology*, Proceedings of the IEEE (Special Issue on Pattern Recognition and Image Processing), 67 (5), 708–713.
- Fisher R. A., (1936), The Use of Multiple Measurements in Taxonomic Problems, *Annals of Eugenics*, 7 (2), 179–188.
- Hadasik D., (1998), *Upadłość przedsiębiorstw w Polsce i metody jej prognozowania*, Zeszyty Naukowe – seria II, Prace habilitacyjne, Zeszyt 153, Akademia Ekonomiczna w Poznaniu, Poznań.
- Hołda A., (2000), *Optymalizacja i model zastosowania procedur analitycznych w rewizji sprawozdań finansowych*, praca doktorska, Akademia Ekonomiczna w Krakowie, Kraków.
- Hołda A., (2001), Prognozowanie bankructwa jednostki w warunkach gospodarki polskiej z wykorzystaniem funkcji dyskryminacyjnej ZH, *Rachunkowość*, 5, 306–310.
- Hołda A., (2006), *Zasada kontynuacji działalności i prognozowanie upadłości w polskich realiach gospodarczych*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.
- Korol T., (2010), *Systemy ostrzegania przedsiębiorstw przed ryzykiem upadłości*, Wolters Kluwer, Warszawa.
- Korol T., Prusak B., (2005), *Upadłość przedsiębiorstw a wykorzystanie sztucznej inteligencji*, CeDeWu, Warszawa.
- Krzyśko M., Wołyński W., Górecki T., Skorzybut M., (2008), *Systemy uczące się. Rozpoznawanie wzorców, analiza skupień i redukcja wymiarowości*, Wydawnictwo Naukowo-Techniczne w Warszawie, Warszawa.
- Kuncheva L. I., (2000), *Cluster-and-Selection Method for Classifier Combination*, 4th International Conference on Knowledge-Based Intelligent Engineering Systems & Allied Technologies, Brighton, UK.
- Kuncheva L. I., (2004), *Combining Pattern Classifiers: Methods and Algorithms*, Wiley.
- Mączyńska E., (1994), Ocena kondycji przedsiębiorstwa (Uproszczone metody), *Życie gospodarcze*, 38, 42–45.
- Mączyńska E., Zawadzki M., (2006), Dyskryminacyjne modele predykcji bankructwa przedsiębiorstw, *Ekonomista*, 2, 205–235.
- Mossman C. E., Bell G. G., Swartz L. M., Turtle H., (1998), An Empirical Comparison of Bankruptcy Models, *The Financial Review*, 33 (2), 35–54.
- Ohlson J. A., (1980), Financial Ratios and the Probabilistic Prediction of Bankruptcy, *Journal of Accounting Research*, 18 (1), 109–131.
- Pociecha J., Pawełek B., Baryła M., Augustyn S., (2014), *Statystyczne metody prognozowania bankructwa w zmieniającej się koniunkturze gospodarczej*, Fundacja Uniwersytetu Ekonomicznego w Krakowie, Kraków.
- Wędzki D., (2005), Bankruptcy Logit Model for Polish Economy, *Argumenta Oeconomica Cracoviensia*, 3, 49–70.
- Woods K., Kegelmeyer W. P., Bowyer K., (1997), Combination of Multiple Classifiers Using Local Accuracy Estimates, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 19 (4), 405–410.

METODY ŁĄCZENIA I SELEKCJI KLASYFIKATORÓW W PROGNOZOWANIU UPADŁOŚCI PRZEDSIĘBIORSTW

Streszczenie

Autor niniejszego artykułu postanowił zbadać skuteczność wykorzystania metod łączenia oraz selekcji klasyfikatorów w prognozowaniu upadłości przedsiębiorstw w Polsce. Przeprowadzone badanie pozwoliło na porównanie jakości stawianych prognoz przez cztery klasyfikatory indywidualne: liniowa analiza dyskryminacyjna, regresja logistyczna, sztuczna sieć neuronowa oraz las losowy z wynikami

dziewięciu metod łączenia oraz selekcji, bazujących na zbiorze powyższych klasyfikatorów. Autor artykułu przeprowadził także analizę wpływu liczby uwzględnianych zmiennych na poprawność klasyfikacji poszczególnych metod.

Słowa kluczowe: prognozowanie upadłości przedsiębiorstw, metody klasyfikacyjne, metody łączenia i selekcji klasyfikatorów

CLASSIFIER COMBINATION AND SELECTION METHODS AT THE CORPORATE BANKRUPTCY PREDICTION

A b s t r a c t

The author of this article decided to investigate the efficiency of using classifier combination and selection methods at the prediction of corporate bankruptcy in Poland. The study allowed us to compare the predictions from 4 individual classifiers: linear discriminant analysis, logistic regression, artificial neural network and random forest with the predictions from 9 classifier combination and selection methods. Moreover, the author analyzed the influence of the number of variables in models on performances of the correct classification for all considered methods.

Keywords: corporate bankruptcy prediction, classification methods, classifier combination and selection methods

GRZEGORZ TARCZYŃSKI¹MODEL DLA ZADANIA KOMPLETACJI ZLECEŃ ŁĄCZONYCH
UWZGLĘDNIAJĄCY PROBLEM BLOKOWANIA SIĘ MAGAZYNIERÓW

1. WSTĘP

Dynamiczny rozwój handlu internetowego pociąga za sobą konieczność szybkiej i niezawodnej kompletacji zamówień. Klienci bowiem często już następnego dnia chcą otrzymać zamówiony towar. Z tego powodu wdrażane są w magazynach nowoczesne rozwiązania mające umożliwić bardziej płynną realizację procesu kompletacji. Jednym z takich rozwiązań jest instalacja systemu automatycznego, w którym towary (umieszczone zazwyczaj w kuwetach) za pomocą suwnicy i systemu przenośników dostarczane są do stanowiska kompletacyjnego. Zakup takiego systemu generuje jednak poważne koszty, dlatego wciąż powszechnie stosowana jest kompletacja tradycyjna, odbywająca się zgodnie z zasadą „człowiek do towaru”. Również w takim przypadku możliwe jest znaczne zwiększenie efektywności (rozumianej jako minimalizacja czasu, jaki magazynier przeznacza na skompletowanie zamówienia) procesu realizacji zamówień. Takim sposobem jest m.in. równoczesna kompletacja kilku zamówień poprzez stworzenie tzw. zleceń łączonych. W ten sposób średni dystans przypadający na jedno zamówienie, a jaki jest pokonywany przez magazyniera, ulega skróceniu, co z kolei skutkuje redukcją czasów kompletacji.

Problem badania wpływu tworzenia zleceń łączonych na średnie czasy kompletacji zamówień wydaje się nie do końca zbadany. Magazynierzy, którzy pobierają w jednym cyklu większą liczbę towarów, mogą przeszkadzać sobie w pracy – zwłaszcza w sytuacji, gdy w magazynie towary składowane są w oparciu o klasyfikację ABC. Umieszczenie towarów szybko rotujących w takich rejonach magazynu, które są łatwo dostępne, jest bowiem kolejnym elementem mogącym spowodować redukcję pokonywanego podczas kompletacji dystansu. Jednym ze sposobów, stosowanych do podziału magazynu na strefy ABC jest polityka *within-aisle*, gdzie klasę A, w której umieszcza się towary najszybciej rotujące, tworzy pewna liczba alejek najbliższych miejsca, w którym magazynier rozpoczyna i kończy swoją pracę. Na kolejne klasy – zawierające towary coraz wolniej rotujące – przeznacza się alejki coraz bardziej oddalone. W ramach każdej klasy towary rozmieszcza się w sposób losowy, a w bada-

¹ Uniwersytet Ekonomiczny we Wrocławiu, Wydział Zarządzania, Informatyki i Finansów, Katedra Ekonometrii i Badań Operacyjnych, ul. Komandorska 118/120, 53-345 Wrocław, Polska, e-mail: grzegorz.tarczyński@ue.wroc.pl.

niach teoretycznych przyjmuje się dodatkowe założenie o braku korelacji pomiędzy towarami (pojawienie się na zamówieniu określonego towaru nie wpływa na prawdopodobieństwo wystąpienia na tym samym zamówieniu innego towaru). Warunek komplementarności towarów uwzględniany jest zazwyczaj podczas analizy tzw. składowania dedykowanego.

W literaturze procesy magazynowe modelowane są zarówno za pomocą narzędzi symulacyjnych, jak i podejścia analitycznego. Symulacje umożliwiają uwzględnienie większej liczby parametrów, ich wadą jest jednak długi czas obliczeń. Celem artykułu jest stworzenie modelu analitycznego umożliwiającego ustalenie optymalnej liczby łączonych zamówień dla dwóch heurystyk wyznaczania trasy magazyniera, które mają największe znaczenie praktyczne. Ponieważ podczas równoczesnej pracy wielu magazynierów w magazynie mogą powstawać zatory, niezbędne jest uwzględnienie w proponowanym modelu ewentualnych interakcji pomiędzy pracownikami. Zgodnie z wiedzą autora, takie badania przeprowadzane były jak dotąd za pomocą narzędzi symulacyjnych. W tej pracy połączono podejście analityczne z symulacyjnym – symulacje wykorzystano do stworzenia postaci analitycznej funkcji czasu blokowania się magazynierów.

Analizując procesy magazynowe, w badaniach teoretycznych zazwyczaj zakłada się składowanie tego samego towaru wyłącznie w jednym miejscu w magazynie. Problem optymalizacji wyboru lokalizacji, z której należy pobrać towar jest więc pomijany w badaniach. Takie założenie zostało przyjęte również przez autora artykułu. Studia nad tym zagadnieniem przedstawione są m.in. przez Dmytrowa (2013, 2015) oraz Dmytrowa, Doszyna (2015).

Artykuł składa się ze wstępu, pięciu rozdziałów oraz wniosków końcowych. W rozdziale 2 przedstawiono przegląd literatury dotyczącej kompletacji zleceń łączonych oraz problemu blokowania się magazynierów. Lista założeń odnoszących się do opracowanego modelu wraz z przyjętą notacją podana jest w rozdziale 3. W rozdziale 4 – na podstawie wyników eksperymentów symulacyjnych – przedstawiono postać analityczną funkcji czasu blokowania. Kolejny rozdział zawiera model pozwalający na optymalizację procesu tworzenia zleceń łączonych. Przykład empiryczny pokazano w rozdziale 6.

2. PRZEGLĄD LITERATURY

Kompletacja zleceń łączonych oprócz oczywistych korzyści, takich jak bardziej efektywne wykorzystanie urządzeń i czasu pracy, generuje też pewne zagrożenia: np. brak możliwości pobrania zbyt dużej liczby towarów w jednym cyklu (Krawczyk, 2011). Innym problemem, który pojawia się podczas kompletacji zleceń łączonych jest ryzyko zwiększonego blokowania się magazynierów (Parikh, Meller, 2008). Problem wzajemnych interakcji zachodzących pomiędzy magazynierami nie został jak dotąd zbadany w wyczerpujący sposób. Zagadnienie to najczęściej analizowane jest za pomocą symulacji, choć powstają również modele analityczne pozwalające na

określenie, w jakich warunkach magazynierzy w najmniejszy oraz największy sposób przeszkadzają sobie w pracy. Dla magazynów z wąskimi alejkami, w których można w precyzyjny i jednoznaczny sposób wyznaczyć niezmienną trasę, po której poruszać się będą magazynierzy, sama kompletacja towarów, jak i blokowanie się magazynierów, mogą być analizowane za pomocą procesu Markowa (Gue i inni, 2006; Hong i inni, 2010; Parikh, Meller, 2010). W pracach tych zakłada się stałą wartość prawdopodobieństwa pobrania towarów dla wszystkich lokalizacji w magazynie i w sposób analityczny badane są dwa przypadki: taki w którym magazynierzy podczas procesu kompletacji zamówień przemieszczają się po alejkach bardzo wolno oraz taki, w którym poruszają się oni nieskończenie szybko – stosunek czasu pobrania towaru z określonej lokalizacji do czasu przejścia obok lokalizacji z towarem (ang. *pick:walk ratio*) wynosi odpowiednio 1:1 i ∞ :1. Inne wartości tej relacji analizowane są z wykorzystaniem symulacji. Dopiero Hong (2014) wykorzystał proces Markowa do stworzenia bardziej ogólnego modelu $m:1$ dla dwóch magazynierów. Gue i inni (2006) pokazali, że dla wariantu 1:1 problem blokowania występuje w największym stopniu wtedy, gdy magazynierzy muszą pobrać towary średnio z około co trzeciej lokalizacji. Dla innych wariantów zatory w magazynie pojawiały się najrzadziej, gdy prawdopodobieństwa pobrania towarów były bardzo niskie, albo bardzo wysokie. Parikh, Meller (2010) zbadali wariant, w którym czasy pobrania towarów były losowe. O ile dla wariantu 1:1 i przypadku deterministycznego problem blokowania nasilał się, gdy prawdopodobieństwo pobrania towarów z lokalizacji wynosiło około 35%, to dla przypadku stochastycznego największe zatory pojawiały się, gdy magazynierzy pobierali towary ze wszystkich lokalizacji w magazynie (prawdopodobieństwo równe 100%). Podobne badanie wykorzystujące proces Markowa, ale dla magazynów z szerokimi alejkami przeprowadzili Parikh, Meller (2009).

Zagadnienie kompletacji towarów i blokowania magazynierów modelowane jest również za pomocą systemów kolejkowych. Pan i inni (2012) zbadali magazyny z alejkami o ruchu dwustronnym przy jednoczesnym założeniu, że w alejce może w tym samym czasie przebywać tylko jedna osoba. Zhang i inni (2009) analizowali 4 przypadki, które mogą prowadzić do powstawania zatorów w drukarni. Autorzy zalecają, aby zarówno w sytuacjach, gdy natężenie ruchu jest niewielkie, jak i wtedy, gdy jest bardzo wysokie, poruszać się zawsze po najkrótszych trasach. Poszukiwanie dróg alternatywnych przynosi efekty, gdy intensywność ruchu jest na średnim poziomie.

Chen i inni (2016) dla magazynu wieloblokowego oraz losowych czasów pobrania towarów zaproponowali, aby w celu uniknięcia blokowania modyfikować trasy magazynierów z wykorzystaniem algorytmów mrówkowych na bieżąco – w trakcie realizacji procesu kompletacji.

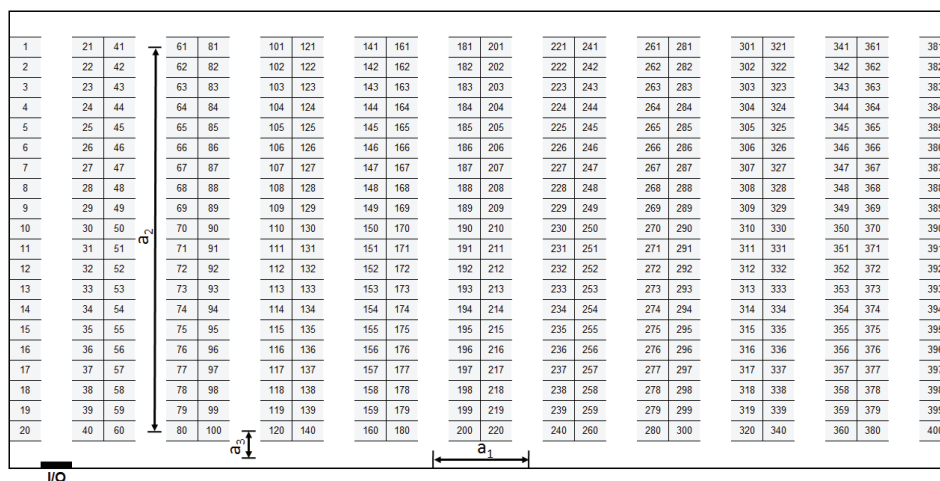
Hong i inni (2012a) przedstawili model oraz zaproponowali heurystykę pozwalającą na łączenie zamówień z rozważaniem możliwości blokowania się magazynierów. Autorzy zaprezentowali metodę ustalania kolejności realizacji zleceń, która minimalizuje czasy kompletacji (z uwzględnieniem problemu zatorów) przy założeniu deterministycznych czasów pobrania towarów. Hong i inni (2012b) problem łączenia

zamówień potraktowali w oryginalny sposób: wyznaczyli wszystkie możliwe trasy przejścia przez prostokątny jednoblokowy magazyn z wąskimi alejkami, a następnie do tych tras przypisywali zamówienia – w ten sposób scalając je ze sobą. Przedstawiony model minimalizuje odległość pokonywaną przez magazynierów, a ewentualne zatory zbadane zostały za pomocą symulacji.

Parikh, Meller (2008) w celu rozwiązania zadania łączenia zamówień poszukiwali analogii do problemu pakowania (ang. *bin packing problem*). Autorzy zaprezentowali zmodyfikowany dualny problem pakowania: wyznaczenie maksymalnej liczby towarów możliwych do pobrania przez zadaną liczbę magazynierów.

3. ZAŁOŻENIA DO MODELU POZWALAJĄCEGO NA ŁĄCZENIE ZAMÓWIEŃ

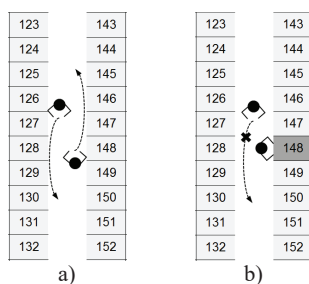
W badaniach analizie poddany będzie najczęściej opisywany w literaturze magazyn prostokątny jednoblokowy (rysunek 1), gdzie punkt I/O w którym magazynierzy rozpoczynają i kończą cykl kompletacyjny znajduje się naprzeciwko wejścia do pierwszej alejki. W magazynie znajdują się dwa główne korytarze, umownie zwane dolnym i górnym. Towary składowane są w bocznych alejkach – w badaniach empirycznych przeanalizowano magazyny z 4 i 10 alejkami. Alejki są na tyle szerokie, że magazynierzy mogą się w nich poruszać w dwie strony i w razie potrzeby mijać – za wyjątkiem sytuacji, kiedy jeden z magazynierów pobiera towar (rysunek 2). Wówczas powstaje zator i magazynier chcący przejść obok musi poczekać, aż inny pracownik skończy pobierać towar (blokowanie w szerokich alejkach przedstawione jest np. w artykule Parikha, Mellera, 2009).



Na rysunku zaznaczono odległości: pomiędzy wejściami do dwóch sąsiadujących alejek (a_1), z głównego korytarza do najbliższej lokalizacji w alejce (a_2) i od pierwszej do ostatniej lokalizacji w regale (a_3).

Rysunek 1. Magazyn prostokątny jednoblokowy z 10 alejkami

Źródło: opracowanie własne.



- a) magazynierzy mogą się mijać (zator nie powstaje),
 b) jeden magazynier pobiera towar, a drugi chce przejść (powstaje zator)

Rysunek 2. Powstawanie zatorów w magazynach z szerokim alejkami

Źródło: opracowanie własne na podstawie Parikh, Meller, 2009.

W literaturze opisanych jest 5 heurystyk służących do wyznaczania trasy magazynierów w magazynach prostokątnych jednoblokowych, znany jest też algorytm, oparty na programowaniu dynamicznym, umożliwiający szybkie znalezienie drogi optymalnej (Ratliff, Rosenthal, 1983). Średnie czasy kompletacji dla każdej z wymienionych metod zależą od pokonywanego przez pracowników dystansu, na który wpływa wiele parametrów, takich jak: liczba i długość alejek, lokalizacja punktu I/O, liczba towarów na zamówieniu, sposób rozmieszczenia w magazynie towarów szybko rotujących. Wymienione parametry różnicują dystans, a więc i czasy kompletacji dla różnych sposobów wyznaczania trasy. Inne, jak np. czas inicjalizacji zamówienia, liczba pobieranych towarów tego samego typu i (w przypadku składowania jednopoziomowego) czas pobrania towaru, wpływają na czas kompletacji zamówień, ale zazwyczaj nie zaburzają hierarchii rankingu wyników (choć mogą wpływać na zatory). Czasy kompletacji zamówień dla różnych parametrów mogą być wyznaczane w sposób dokładny z wykorzystaniem symulacji komputerowych i w sposób zazwyczaj przybliżony za pomocą podejścia analitycznego.

W artykule omówione zostaną dwie najpopularniejsze heurystyki: *s-shape* (zwana również *traversal* lub *transversal*) oraz *return*, dla których można w dość dokładny sposób wyznaczać średnie czasy kompletacji zamówień w sposób analityczny (np. Tarczyński, 2015a, 2015b). Metody te są łatwo implementowalne i mają duże znaczenie praktyczne, w odróżnieniu od pozostałych opisanych w literaturze: zarówno heurystyk, jak i algorytmu drogi optymalnej (ze względu na trudności z wdrożeniem wykorzystywane są raczej tylko w badaniach teoretycznych). Trasa wyznaczona za pomocą metody *s-shape* kształtem przypomina literę „S”. Magazynier wchodzi tylko do tych alejek, w których składowane są towary, które znajdują się na aktualnie obsługiwanym zleceniu (liście kompletacyjnej). Po wejściu do alejki pobierane są wszystkie potrzebne towary, po czym pracownik opuszcza alejkę drugim wejściem. Podobne zasady obowiązują dla heurystyki *return*, z tą różnicą, że tutaj w użyciu jest tylko dolny główny korytarz, a więc po pobraniu z alejki wyróżnionych na zleceniu towarów magazynier zawsze musi zawrócić.

W dalszych obliczeniach przyjęto następującą notację:

- i, j – numer alejki, $i, j = 1 \dots K$,
- K – liczba alejek w magazynie,
- n – liczba towarów na zamówieniu,
- n_i – średnia liczba towarów pobieranych przez magazynierów z i -tej alejki,
- m – liczba magazynierów pracujących równocześnie w magazynie,
- m_i – średnia liczba magazynierów pracujących równocześnie w i -tej alejce,
- B_i – średni czas blokowania w i -tej alejce,
- a_1 – odległość pomiędzy wejściami do dwóch sąsiadujących alejek,
- a_2 – długość bocznych alejek (odległość od pierwszej do ostatniej lokalizacji),
- a_3 – odległość od środka głównego korytarza do pierwszego miejsca składowania towaru w alejce,
- t_{pobr} – czas pobrania towaru,
- t_{ini} – suma czasów inicjalizacji i zakończenia zamówienia,
- t_{bl} – czas stracony przez blokowanie podczas kompletacji jednego zamówienia,
- t_{sort} – czas sortowania przypadający na 1 towar,
- p_i – prawdopodobieństwo, że towar z zamówienia znajduje się w i -tej alejce,
- $\dot{P}_i$ – prawdopodobieństwo wejścia przez magazyniera do i -tej bocznej alejki

$$\dot{P}_i = 1 - (1 - p_i)^n,$$

- $\ddot{P}_i$ – prawdopodobieństwo dojścia przez magazyniera do wejścia do i -tej bocznej alejki

$$\ddot{P}_i = 1 - \left(1 - \sum_{j=i}^K p_j \right)^n,$$

- d_1 – średnia odległość pokonywana przez magazyniera w głównych korytarzach,
- d_2^{heur} – średnia odległość pokonywana przez magazyniera (poruszającego się zgodnie z heurystyką *heur*) w bocznych alejkach,
- t^{heur} – średni czas kompletacji zamówienia dla heurystyki *heur*,
- v – prędkość z jaką porusza się magazynier,
- x – liczba łączonych zamówień,
- $f(x)$ – funkcja wyrażająca procentową stratę czasową na pobranie towaru wynikającą z połączenia zamówień.

4. ŚREDNI CZAS BLOKOWANIA

W celu wyznaczenia postaci analitycznej funkcji czasu blokowania podczas procesu kompletacji zamówień wykonano szereg symulacji z wykorzystaniem programu Warehouse Real-Time Simulator (Tarczyński, 2013). Istotnym elementem

oceny jakości modelu symulacyjnego jest jego weryfikacja, czyli „sprawdzenie czy model komputerowy jest wystarczająco trafną reprezentacją modelu konceptualnego” (Schlesinger i inni, 1979, cyt. za Balcerak, 2003). Wyniki uzyskiwane przez program Warehouse Real-Time Simulator zostały zweryfikowane na różnych poziomach szczegółowości. Dla każdego zrealizowanego zamówienia program generuje statystyki dotyczące m.in. numerów odwiedzanych lokalizacji, momentu rozpoczęcia i zakończenia procesu kompletacji, ale i dystansu pokonanego przez magazyniera zarówno w głównych korytarzach, jak i w bocznych alejkach. Wyniki te zostały sprawdzone dla różnych parametrów wejściowych z wynikiem pozytywnym. Kolejnym etapem weryfikacji wyników uzyskanych za pomocą symulacji jest porównanie średnich czasów kompletacji zamówień z modelami analitycznymi – dla heurystyk *return* i *s-shape* różnica uzyskanych wyników nie przekracza 1% dla dowolnych prawdopodobieństw pobierania towarów z poszczególnych lokalizacji magazynowych (autor przygotowuje artykuł na ten temat). O ile łatwo zweryfikować jakość modelu symulacyjnego, gdy w magazynie pracuje jedna osoba, o tyle w przypadku badania interakcji pomiędzy pracownikami nie jest to już takie proste. Sytuacje blokowania się magazynierów można jednak sprawdzić za pomocą analizy wizualnej (program Warehouse Real-Time Simulator przewiduje możliwość śledzenia realizacji procesu kompletacji zamówień za pomocą animacji komputerowej). Istniejące modele analityczne wymagają przyjęcia pewnych założeń, które nie znajdują odzwierciedlenia w proponowanym modelu, nie pozwalają więc na zweryfikowanie uzyskanych wyników. Wyniki te należałoby poddać walidacji, czyli sprawdzeniu w jakim stopniu model jest wiernym odwzorowaniem rzeczywistości (Balcerak, 2003; Sargent, 2013). Zadanie to jest trudne do wykonania, w modelu odzwierciedlono jednak sytuacje blokowania opisane przez Parikha, Meller (2009). Przedstawienie postaci analitycznej funkcji czasu blokowania magazynierów stanowi jeden z celów artykułu. Jest ona niezbędna do wyznaczenia optymalnej liczby łączonych zleceń, co stanowi główny cel pracy.

Do badań przyjęto wartość współczynnika *pick:walk* na poziomie 30:1. Zbadano jak na średni czas blokowania w jednej alejce wpływa liczba pracowników jednocześnie w niej przebywających oraz liczba towarów pobieranych przez magazynierów. Symulacje zrealizowano dla liczby pracowników równej od 1 do 10 i liczby towarów od 1 do 15. Analizie poddane zostały dwie heurystyki służące do wyznaczania trasy magazyniera: *s-shape* i *return*. Dla każdego eksperymentu wykonano 20000 replikacji. De Koster i inni (1998) pokazali, że przy założeniu, że czasy kompletacji zamówień mają rozkład normalny liczba 10000 replikacji jest zadowalająca. W przeprowadzonych eksperymentach badano rozkłady czasów blokowania i otrzymano maksymalne wartości kurtozy (wszystkie są dodatnie): 1,16 dla heurystyki *return* i 1,83 dla *s-shape* oraz współczynnika skośności (wszystkie rozkłady są prawostronnie asymetryczne): odpowiednio 0,13 i 0,19. Przy założeniu normalności rozkładów, dla 99% poziomu ufności i maksymalnego błędu szacunku równego 1 sekundzie – otrzymano minimalne liczby replikacji. Dla metody *s-shape* wartość ta nie przekraczała 15031, a dla

return 18053. Liczba replikacji na poziomie 20000 jest więc wystarczająca do uzyskania dokładnych wyników we wszystkich eksperymentach. Założenia dotyczące symulacji przedstawiono w tabeli 1.

Tabela 1.

Założenia dotyczące symulacji komputerowej

Zmienna wynikowa	
Czas blokowania magazyniera	
Parametry	Dziedzina
Liczba symulowanych alejek	1
Współczynnik <i>pick:walk</i>	30:1
Liczba pracowników	1–10
Liczba towarów pobieranych w jednej alejce	1–15
Heurystyka wyznaczania trasy	<i>s-shape, return</i>
Liczba eksperymentów	10 x 15 x 2 = 300
Parametry wpływające na jakość uzyskanych wyników	Zadana wartość
Poziom ufności	99%
Maksymalny błąd szacunku	1 sekunda
Liczba replikacji dla każdego eksperymentu	20000

Źródło: opracowanie własne.

Z przeprowadzonych badań wynika, że średni czas blokowania magazyniera w alejce zależy liniowo od liczby przebywających w niej pracowników i nieliniowo od liczby pobieranych towarów. Po wizualnej analizie wykresów, do wyznaczania czasu blokowania w alejce zaproponowano postać analityczną funkcji:

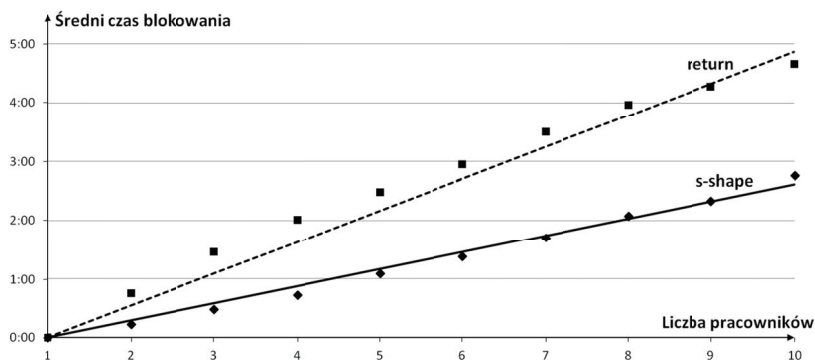
$$B_i(m_i, n_i) = \alpha(m_i - 1)\tanh(\beta n_i + \gamma), \quad (1)$$

gdzie: α, β, γ – parametry podlegające estymacji.

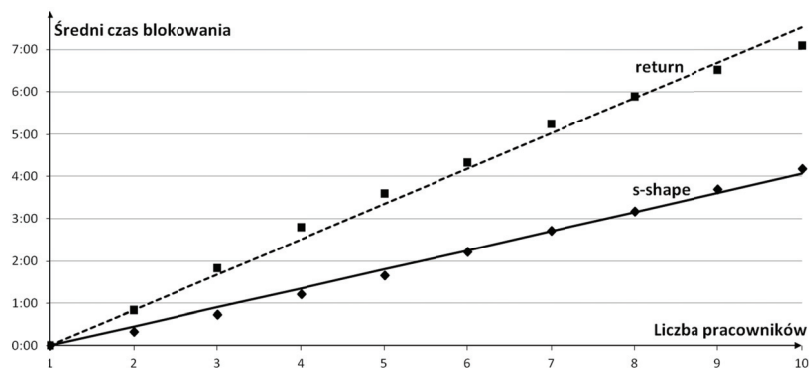
Z wykorzystaniem metody najmniejszych kwadratów oszacowano wartości parametrów dla różnych metod wyznaczania trasy:

- *s-shape*: $\alpha = 0,92$, $\beta = 0,09$, $\gamma = 0,11$, $R^2 = 0,993$,
- *return*: $\alpha = 1,5$, $\beta = 0,12$, $\gamma = 0,07$, $R^2 = 0,995$.

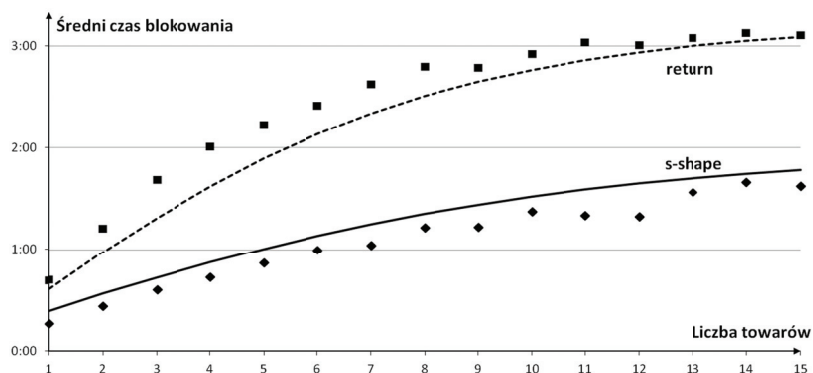
Przykładowe wykresy ilustrujące zależność średniego czasu blokowania od liczby pracowników oraz liczby pobieranych towarów przedstawione są na rysunkach 3–6. Punkty przedstawiają wartości uzyskane za pomocą symulacji, natomiast krzywe uzyskano w wyniku zastosowania wzoru (1).



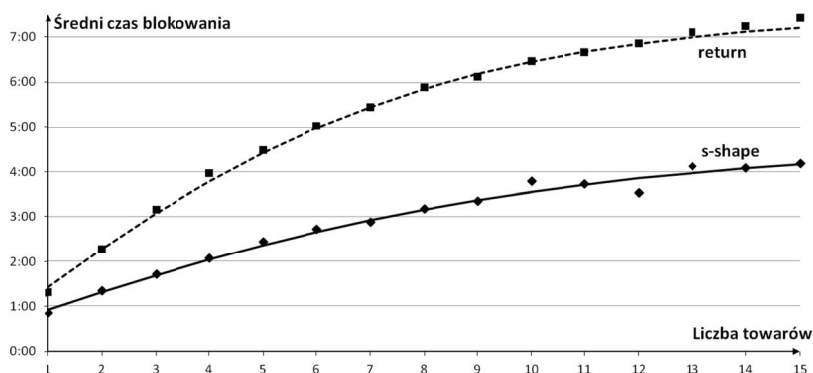
Rysunek 3. Zależność czasu blokowania od liczby pracowników dla 4 towarów pobieranych w alejce
Źródło: opracowanie własne.



Rysunek 4. Zależność czasu blokowania od liczby pracowników dla 8 towarów pobieranych w alejce
Źródło: opracowanie własne.



Rysunek 5. Zależność czasu blokowania od liczby towarów dla 4 pracowników przebywających w alejce
Źródło: opracowanie własne.



Rysunek 6. Zależność czasu blokowania od liczby towarów dla 8 pracowników przebywających w alejce

Źródło: opracowanie własne.

5. MODEL ŁĄCZĄCY ZAMÓWIENIA MINIMALIZUJĄCY ŚREDNI CZAS KOMPLETACJI ZAMÓWIEŃ

Na odległość pokonywaną przez magazynierów podczas kompletacji wyrobów składa się dystans przemierzany w głównych korytarzach i bocznych alejkach. Odległość, którą pracownicy muszą przebyć w głównych korytarzach nie zależy od przyjętej heurystyki wyznaczania trasy i w przypadku całkowicie losowego rozmieszczenia towarów w alejkach może być wyliczona zgodnie ze wzorem podanym np. w pracy Tarczyńskiego (2015a). Propozycja ta wymaga drobnych korekt, uwzględniających możliwość rozmieszczenia w magazynie towarów szybko rotujących również w sposób losowy, ale w osobnych alejkach (zgodnie z klasyfikacją ABC dla towarów z każdej klasy wyróżnione są całe alejki).

Wzór można wyprowadzić korzystając z własności rozkładu dwumianowego. Zamówienia generowane są w sposób losowy. Popyt na poszczególne towary jest zróżnicowany, w praktyce zazwyczaj towary uszeregowane malejąco względem popytu spełniają regułę Pareto (za 80% skutków odpowiada 20% przyczyn) i popyt na nie może być modelowany np. za pomocą rozkładu Pareto. Niech wartości p_j oznaczają skumulowane wartości prawdopodobieństwa dla wszystkich towarów przechowywanych w j -tej alejce. Wyrażenie $1 - \sum_{j=i}^K p_j$ stanowi prawdopodobieństwo, że towar z zamówienia znajduje się w alejkach oznaczonych numerami mniejszymi od i . Na zamówieniu znajduje się n towarów (zakładamy, że z każdej alejki znajduje się tak dużo towarów, że do tworzenia zamówień można wykorzystać metodę losowania z powtórzeniami). Prawdopodobieństwo, że wszystkie towary z zamówienia znajdują się w alejkach o numerach mniejszych od i wynosi $(1 - \sum_{j=i}^K p_j)^n$. Magazynier dojdzie do i -tej alejki tylko wtedy, gdy będzie miał do pobrania chociaż jeden towar z alejek o numerach co najmniej i , a więc wszystkie towary nie mogą znajdować się w alejkach o numerach mniejszych od i : stąd prawdopodobieństwo dojścia do i -tej alejki wynosi

$1 - (1 - \sum_{j=i}^K p_j)^n$. Sumując prawdopodobieństwa dojścia do wszystkich alejek i mnożąc wynik przez odległość pomiędzy wejściami do dwóch sąsiadujących alejek otrzymujemy wartość oczekiwaną odległości pokonywanej przez magazyniera w głównych korytarzach (wynik należy jeszcze przemnożyć przez 2, ponieważ dystans w głównych korytarzach pokonywany jest dwukrotnie, a od sumy prawdopodobieństw odjąć 1, bo prawdopodobieństwo znalezienia się magazyniera w pierwszej alejce wynosi 1, a przyjmuje się, że odległość z punktu I/O, gdzie magazynier zaczyna i kończy pracę do wejścia do pierwszej alejki jest nieistotna). Średnia odległość pokonywana w głównych korytarzach wynosi więc:

$$d_1 = 2 \left(\sum_{i=1}^K \dot{P}_i - 1 \right) a_1 = 2 \left(\sum_{i=1}^K \left(1 - \left(1 - \sum_{j=i}^K p_j \right)^n \right) - 1 \right) a_1. \quad (2)$$

W przypadku bocznych alejek, dla heurystyki *s-shape* w celu obliczenia przybliżonej wartości przebytego dystansu należy wyznaczyć wartość oczekiwaną liczby alejek, do których wejdzie magazynier ($\sum_{i=1}^K \dot{P}_i$) i przemnożyć ją przez długość alejek. Dla metody *return* trzeba wyznaczyć wartość oczekiwaną liczby towarów pobieranych w *i*-tej alejce pod warunkiem, że magazynier wejdzie do tej alejki: ($\frac{np_i}{\dot{P}_i}$) i do wyznaczenia średniej pozycji najbardziej odległego towaru, który należy pobrać w tej alejce skorzystać ze znanego wzoru na wartość oczekiwaną maksimum podczas losowania *n* liczb z rozkładu jednostajnego (0,1): $E(\max\{x_1, x_2, \dots, x_r\}) = \frac{r}{r+1}$. Modyfikując wzory podane przez Tarczyńskiego (2015a, 2015b) na potrzeby klasyfikacji ABC otrzymujemy:

$$d_2^{s-shape} = (a_2 + 2a_3) \sum_{i=1}^K \dot{P}_i, \quad (3)$$

$$d_2^{return} = 2a_2 \sum_{i=1}^K \dot{P}_i \left(\frac{\frac{np_i}{\dot{P}_i}}{\frac{np_i}{\dot{P}_i} + 1} \right) + 2a_3 \sum_{i=1}^K \dot{P}_i. \quad (4)$$

Model pozwalający na poszukiwanie optymalnej liczby zamówień *x* podlegających łączeniu, przy którym średni czas kompletacji będzie minimalny przybiera postać:

$$\frac{1}{x} t^{heur}(xn) \rightarrow \min,$$

gdzie: $x \geq 1$.

Średni czas kompletacji zamówień zależy od przyjętej polityki kompletacji zleceń łączonych. Stosując metodę *sort-while-pick* magazynierzy przemieszczają się po magazynie z kilkoma pojemnikami, w których umieszczają towary z poszczególnych zamówień. Przy metodzie *pick-then-sort* wszystkie towary wstępnie pobierane są do jednego pojemnika, a dopiero później są sortowane.

Dla polityki *pick-then-sort*:

$$t^{heur}(n) = \frac{d_1 + d_2^{heur}}{v} + nt_{pobr} + t_{ini} + t_{bl}(n) + nt_{sort}. \quad (5)$$

Dla polityki *sort-while-pick*:

$$t^{heur}(n) = \frac{d_1 + d_2^{heur}}{v} + n(1 + f(x))t_{pobr} + t_{ini} + t_{bl}(n). \quad (6)$$

We wzorze (5) w metodzie *pick-then-sort* do czasu kompletacji zamówień dodano czas sortowania wyrobów. W praktyce sortowanie wykonywane jest zazwyczaj w przez innych pracowników lub system automatycznego sortowania. Wartość t_{sort} jest więc równa zero. Składnik ten dodano do wzoru, aby podkreślić konieczność przeprowadzenia czynności sortowania zamówień, co wpływa na koszty przeprowadzenia procesu kompletacji, a może też implikować potrzebę wydzielenia pracowników do sortowni (co przy założeniu stałej liczby zatrudnionych magazynierów spowoduje, że mniej osób będzie pobierać towary z lokalizacji).

Średni czas blokowania przypadający na kompletację jednego zlecenia, na którym znajduje się n lokalizacji z towarami można wyznaczyć korzystając z rozkładu dwumianowego. Magazynierzy większość czasu spędzają w alejkach, z których pobierają towary. Dla uproszczenia obliczeń przyjęto założenie, że prawdopodobieństwo znalezienia się magazyniera w określonej alejce jest proporcjonalne do przeciętnej liczby towarów pobieranych z tej alejki. Zakłada się również, że rozkład prawdopodobieństwa wejścia do alejek jest dla każdego pracownika taki sam (wszyscy realizują też zlecenia z taką samą liczbą pozycji). Wówczas wzór na średni czas blokowania podczas kompletacji towarów z jednego zlecenia przyjmie postać:

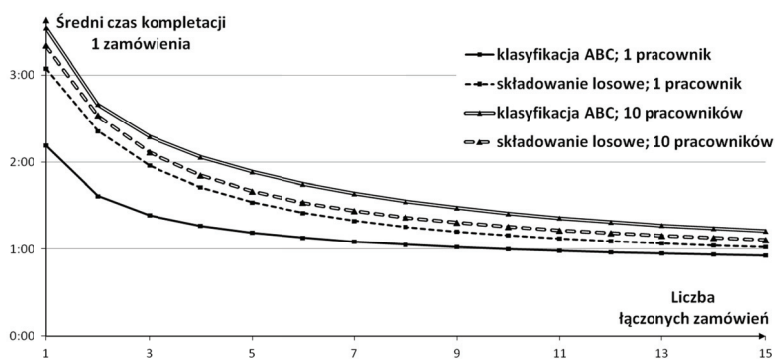
$$t_{bl}(n) = \sum_{i=1}^K p_i \sum_{j=0}^{m-1} \left(\binom{m-1}{j} p_i^j (1 - p_i)^{m-j-1} B_i(j, p_i n) \right). \quad (7)$$

6. PRZYKŁAD EMPIRYCZNY

W badaniach empirycznych sprawdzono jaki wpływ na średnie czasy kompletacji zamówień mają: rozmieszczenie towarów szybko rotujących w alejkach najbliższych punktowi I/O (zgodnie z polityką składowania towarów *within-aisle*, towary w oparciu

o współczynniki rotacji dzieli się na bazie reguły Pareto zazwyczaj na 3 klasy i dla towarów najszybciej rotujących przeznacza się alejki położone najbliższej punktu I/O), wielkość zleceń kompletacyjnych oraz liczba pracujących równocześnie magazynierów. Analizie poddano zamówienia typowe dla firmy internetowej – składające się zazwyczaj z niewielkiej liczby pozycji. Przyjęto, że liczba towarów na jednym zamówieniu jest stała i wynosi 3. Zamówienia przed rozpoczęciem procesu kompletacji łączone były w zlecenia – na jedno zlecenie składało się od 1 do 15 zamówień. W magazynie równocześnie mogło kompletować zamówienia od 1 do 10 osób. Zbadano dwa magazyny: z 4 alejkami i z 10 alejkami, w których trasa kompletacyjna wyznaczana jest wg heurystyki *s-shape* lub *return*.

Średni czas kompletacji zamówień dla magazynów z 4 alejkami przedstawiony jest na rysunkach 7 i 9. W przypadku, gdy zlecenia składają się pojedynczych zamówień zastosowanie składowania opartego na klasyfikacji ABC powoduje oszczędności czasowe od 28,7% do 29,5% w porównaniu ze składowaniem całkowicie losowym. Zwiększając liczbę zamówień kompletowanych w jednym cyklu, korzyści z zastosowania klasyfikacji ABC maleją. W pewnym momencie zbytnie nagromadzenie towarów szybko rotujących na niewielkiej części powierzchni magazynowej zaczyna skutkować tak dużymi zatorami, że taka koncepcja składowania towarów przynosi pogorszenie wyników. Dla heurystyki *s-shape* dzieje się tak od 9 połączonych zamówień, zaś granicą dla heurystyki *return* jest liczba 6 zamówień. Przy zleceniach kompletacyjnych złożonych z 10 zamówień składowanie towarów zgodne z klasyfikacją ABC przynosi pogorszenie wyników o blisko 6% dla heurystyki *s-shape* i 22% dla metody *return* w porównaniu z wariantami, w których zastosowano składowanie całkowicie losowe.

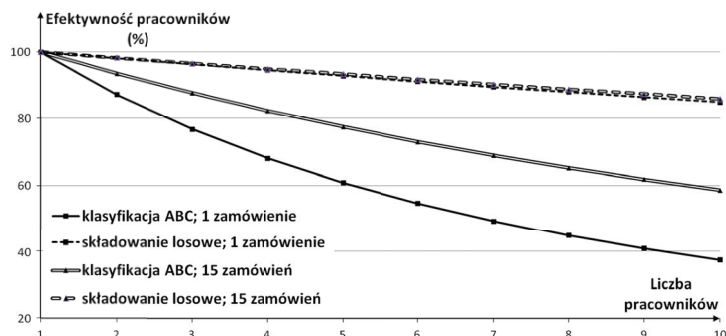


Rysunek 7. Średni czas kompletacji zamówień dla różnej liczby łączonych zamówień i magazynu z 4 alejkami oraz heurystyki *s-shape*

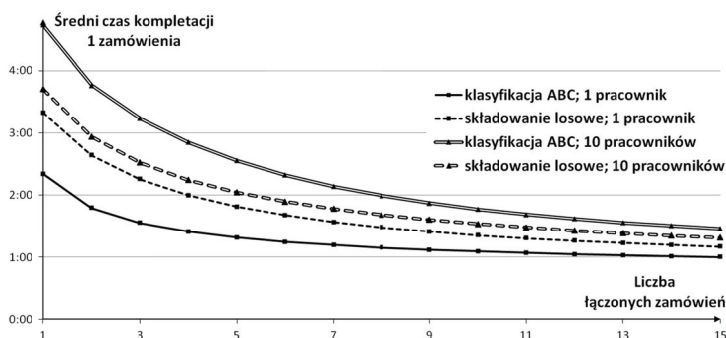
Źródło: opracowanie własne.

Na rysunkach 8 i 10 przedstawiono spadek efektywności kolejnych zatrudnianych w danych warunkach pracowników. Przy losowym składowaniu towarów spadek wydajności pracy jest stosunkowo niewielki: wynosi on dla dziesiątego pracownika około 15% dla heurystyki *s-shape* i 20% dla heurystyki *return* w porównaniu do

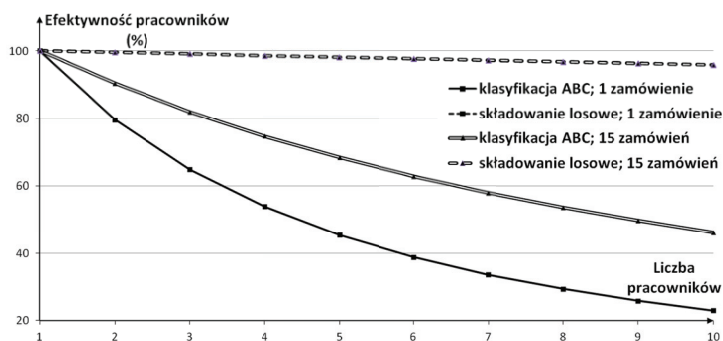
wydajności pierwszego magazyniera. W przypadku składowania towarów zgodnego z klasyfikacją ABC regresja jest dużo silniejsza – dla dziesiątej osoby wynosi ona 41%–63% dla heurystyki *s-shape* i 54%–79% dla heurystyki *return*.



Rysunek 8. Efektywność pracy kolejnych pracowników w magazynie z 4 alejkami dla heurystyki *s-shape*
Źródło: opracowanie własne.

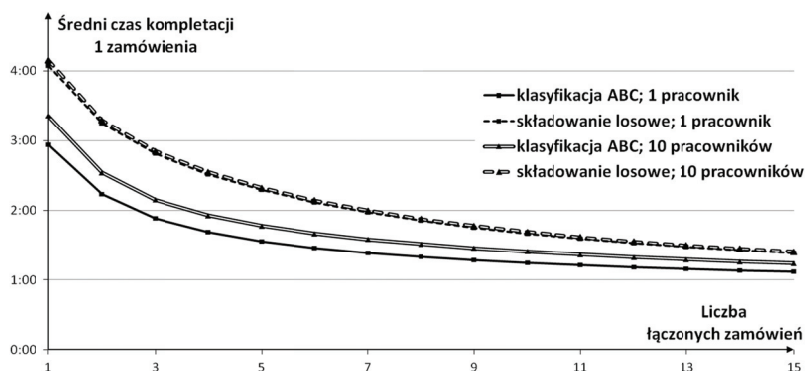


Rysunek 9. Średni czas kompletacji zamówień dla różnej liczby łączonych zamówień i magazynu z 4 alejkami oraz heurystyki *return*
Źródło: opracowanie własne.



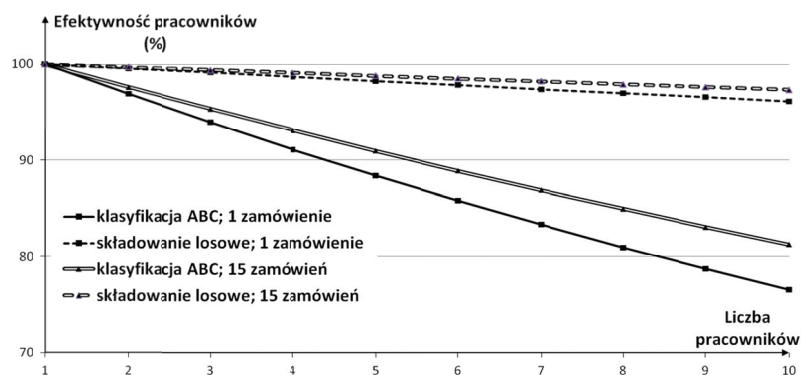
Rysunek 10. Efektywność pracy kolejnych pracowników w magazynie z 4 alejkami dla heurystyki *return*
Źródło: opracowanie własne.

W przypadku magazynów z 10 alejkami, zarówno dla heurystyki *s-shape*, jak i *return*, składowanie towarów zgodne z klasyfikacją ABC zawsze skutkuje redukcją czasu kompletacji zamówień – niezależnie od liczby łączonych zamówień (rysunki 11, 13). Dla metody *s-shape* korzyść wynosi od 11,9% do 33,1%, natomiast dla *return* od 11,0% do 29,7%. Wydajność pracy kolejnych zatrudnionych osób w przypadku składowania losowego spada bardzo nieznacznie – dziesiąty pracownik pracuje mniej efektywnie od pierwszego o około 3% (*s-shape*) i 4% (*return*). W przypadku składowania towarów szybko rotujących w osobnych alejkach regresja wynosi 19%–24% dla heurystyki *s-shape* i 26%–34% dla heurystyki *return*.



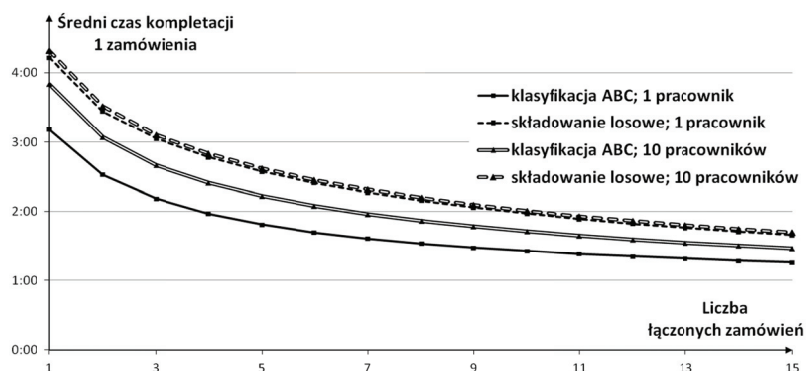
Rysunek 11. Średni czas kompletacji zamówień dla różnej liczby łączonych zamówień i magazynu z 10 alejkami oraz heurystyki *s-shape*

Źródło: opracowanie własne.



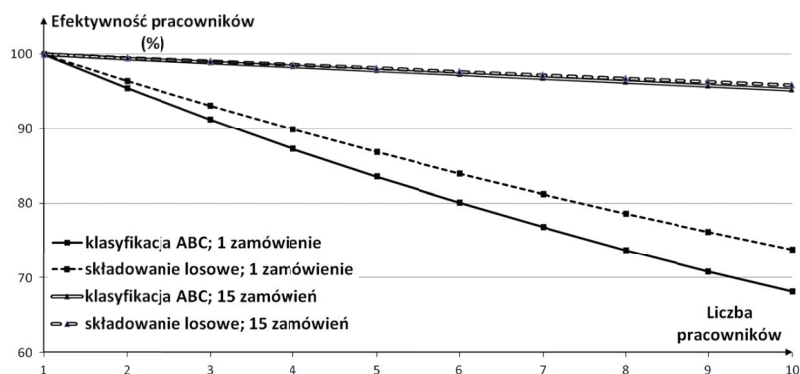
Rysunek 12. Efektywność pracy kolejnych pracowników w magazynie z 10 alejkami dla heurystyki *s-shape*

Źródło: opracowanie własne.



Rysunek 13. Średni czas kompletacji zamówień dla różnej liczby łączonych zamówień i magazynu z 10 alejkami oraz heurystyki *return*

Źródło: opracowanie własne.



Rysunek 14. Efektywność pracy kolejnych pracowników w magazynie z 10 alejkami dla heurystyki *return*

Źródło: opracowanie własne.

7. WNIOSKI

Kompletacja zleceń łączonych powoduje zwiększone ryzyko powstawania w magazynie zatorów. Dla określonej metody wyznaczania trasy i współczynnika *pick:walk* możliwe jest zbudowanie modelu analitycznego wyrażającego czas blokowania magazyniera jako funkcję liczby towarów pobieranych w alejce i liczby przebywających w niej pracowników.

Z wielu badań analizujących wpływ różnych parametrów na czas kompletacji, ale pomijających problem zatorów wynika, że podstawowym sposobem na redukcję czasu kompletacji jest podział towarów ze względu na współczynnik rotacji na klasy ABC i składowanie towarów z każdej klasy w specjalnie wyznaczonych rejonach. W artykule pokazano jednak, że w przypadku kompletacji zleceń łączonych i magazynów o nie-

wielkiej liczbie alejek takie rozwiązanie może nie tylko nie być skuteczne, ale wręcz implikować wydłużeniem czasów kompletacji. Dzieje się tak, ponieważ skoncentrowanie towarów szybko rotujących na małej powierzchni grozi zwiększonym ryzykiem powstawania tam zatorów, które zaburzają płynność procesu. Wydajność strefy kompletacji można podnosić poprzez zwiększanie liczby pracujących w niej magazynierów. Należy jednak pamiętać, że efektywność pracy każdego kolejnego zatrudnionego pracownika spada. Zjawisko to jest szczególnie widoczne w magazynach, w których stosuje się składowanie towarów szybko rotujących w oparciu o politykę *within-aisle*.

Jak dotąd w literaturze brakowało modeli analitycznych optymalizujących liczbę łączonych zamówień, które jednocześnie uwzględniałyby problem blokowania się magazynierów. Badania dotyczące problemu zatorów koncentrowały się na magazynach z wąskimi alejkami z ruchem jednokierunkowym. Przedstawiony w artykule model pozwala na ustalenie liczby łączonych zleceń w magazynach o szerokich alejkach dla dwóch najczęściej stosowanych w praktyce heurystyk wyznaczania trasy magazyniera: *s-shape* i *return*. W modelu uwzględniono dwa typy kompletacji zleceń łączonych: „pobierz i sortuj” oraz „sortuj podczas pobierania”. Dalsze badania powinny dotyczyć wyznaczania liczby, wielkości i kształtów stref, w których składowane są towary o różnych współczynnikach rotacji.

LITERATURA

- Balcerak A., (2003), Walidacja modeli symulacyjnych – źródła postaw badawczych, *Prace Naukowe Instytutu Organizacji i Zarządzania Politechniki Wrocławskiej*, 74 (15), 27–44.
- Chen F., Wang H., Xie Y., Qi C., (2016), An ACO-Based Online Routing Method for Multiple Order Pickers with Congestion Consideration in Warehouse, *Journal of Intelligent Manufacturing*, 27 (2), 389–408.
- De Koster R., Van der Poort E., Roodbergen K. J., (1998), When to Apply Optimal or Heuristic Routing of Orderpickers, *Advances in Distribution Logistics*, Springer Berlin Heidelberg, 375–401.
- Dmytrów K., (2013), Procedura kompletacji zakładająca oczyszczanie lokalizacji, *Zeszyty Naukowe Uniwersytetu Szczecińskiego, Studia i Prace Wydziału Nauk Ekonomicznych i Zarządzania, tom 2 Metody Ilościowe w Ekonomii*, 31, Szczecin, 22–36.
- Dmytrów K., (2015), Taksonomiczne wspomaganie wyboru lokalizacji w procesie kompletacji produktów, *Studia Ekonomiczne*, 248, Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach, Katowice, 17–30.
- Dmytrów K., Doszyń M., (2015), Taksonomiczna procedura wspomaganie kompletacji produktów w magazynie, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu, Taksonomia, 25 Klasyfikacja i analiza danych – teoria i zastosowania*, Wrocław, 72–80.
- Gue K. R., Meller R. D., Skufka J. D., (2006), The Effects of Pick Density on Order Picking Areas with Narrow Aisles, *IIE Transactions*, 38 (10), 859–868.
- Hong S., (2014), Two-worker Blocking Congestion Model with Walk Speed m in a No-passing Circular Passage System, *European Journal of Operational Research*, 235 (3), 687–696.
- Hong S., Johnson A. L., Peters B. A., (2010), Analysis of Picker Blocking in Narrow-Aisle Batch Picking, w: Ellis K., Gue K., Koster R., Meller R., Montreuil B., Oglep M., (red.), *Proceedings of 2010 International Material Handling Research Colloquium (IMHRC)*, The Material Handling Institute, Charlotte, NC, USA.
- Hong S., Johnson A. L., Peters B. A., (2012a), Batch Picking in Narrow-aisle Order Picking Systems with Consideration for Picker Blocking, *European Journal of Operational Research*, 221 (3), 557–570.

- Hong S., Johnson A. L., Peters B. A., (2012b), Large-scale Order Batching in Parallel-aisle Picking Systems, *IIE Transactions*, 44 (2), 88–106.
- Krawczyk S., (red.), (2011), *Logistyka. Teoria i praktyka*, Wydawnictwo Difin, Warszawa.
- Pan J. C. H., Shih P. H., Wu M. H., (2012), Storage Assignment Problem with Travel Distance and Blocking Considerations for a Picker-to-part Order Picking System, *Computers & Industrial Engineering*, 62 (2), 527–535.
- Parikh P. J., Meller R. D., (2008), Selecting Between Batch and Zone Order Picking Strategies in a Distribution Center, *Transportation Research Part E: Logistics and Transportation Review*, 44 (5), 696–719.
- Parikh P. J., Meller R. D., (2009), Estimating Picker Blocking in Wide-aisle Order Picking Systems, *IIE Transactions*, 41 (3), 232–246.
- Parikh P. J., Meller R. D., (2010), A Note on Worker Blocking in Narrow-aisle Order Picking Systems when Pick Time is Non-Deterministic, *IIE Transactions*, 42 (6), 392–404.
- Ratliff H. D., Rosenthal A. S., (1987), Order-picking in a Rectangular Warehouse: A Solvable Case of the Traveling Salesman Problem, *Operation Research*, 31 (3), 515–533.
- Sargent R. G., (2013), Verification and Validation of Simulation Models, *Journal of Simulation*, 7 (1), 12–24.
- Schlesinger S., Crosbie R. E., Gagné R. E., Innis G. S., Lalwani C. S., Loch J., Bartos D., (1979), Terminology for Model Credibility, *Simulation*, 32 (3), 103–104.
- Tarczyński G., (2013), Warehouse Real-Time Simulator – How to Optimize Order Picking Time, Working Paper, Available at SSRN: <http://ssrn.com/abstract=2354827>.
- Tarczyński G., (2015a), Estimating Order-picking Times for Return Heuristic – Equations and Simulations, *LogForum* 11.
- Tarczyński G., (2015b), Średnie czasy kompletacji zamówień dla heurystyki s-shape – wzory i symulacje, *Studia Ekonomiczne 237/15. Informatyka i Ekonometria 2*, Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach, Katowice 2015, 104–116.
- Zhang M., Batta R., Nagi R., (2009), Modeling of Workflow Congestion and Optimization of Flow Routing in a Manufacturing/Warehouse Facility, *Management Science*, 55 (2), 267–280.

MODEL DLA ZADANIA KOMPLETACJI ZLECEŃ ŁĄCZONYCH UWZGLĘDNIAJĄCY PROBLEM BLOKOWANIA SIĘ MAGAZYNIERÓW

Streszczenie

W artykule zaprezentowano model pozwalający na ustalenie optymalnej liczby łączonych zamówień pozwalającej na skrócenie średniego czasu kompletacji. Omówione zostały dwa sposoby kompletacji zleceń łączonych: „pobierz i sortuj” oraz „sortuj podczas pobierania”. Analizie poddano magazyny jednoblokowe prostokątne, gdzie towary składowano w sposób całkowicie losowy oraz losowo w oparciu o klasyfikację ABC. Zbadane zostały dwie heurystyki wyznaczania trasy magazyniera: *s-shape* i *return*, dla których można wyznaczyć wzory na średnie czasy kompletacji korzystając z własności rozkładów prawdopodobieństwa: dwumianowego i jednostajnego. W badaniach uwzględniono problem występowania zatorów w magazynie i na podstawie rezultatów symulacji zaproponowano postać analityczną funkcji czasu blokowania. Z przeprowadzonej analizy wynika, że korzyści czasowe z każdego kolejnego dołączonego do listy kompletacyjnej zamówienia są coraz mniejsze. Wykorzystywanie klasyfikacji ABC do składowania towarów szybko rotujących może w magazynach o niewielkiej liczbie alejek i przy dużych zleceniach kompletacyjnych spowodować wydłużenie czasu kompletacji.

Słowa kluczowe: kompletacja towarów, kompletacja zleceń łączonych, optymalizacja, zatory, blokowanie

MODEL FOR ORDER-BATCHING WITH CONGESTION CONSIDERATION

Abstract

The paper presents the model for optimal number of merged orders that will reduce the average order-picking time. Two order batching policies were described: pick-then-sort and sort-while-pick. One-block rectangular warehouses were considered. The author studied two routing heuristics: s-shape and return for which the equations for average order-picking time can be designated based on probability distributions: binomial distribution and uniform distribution. The research takes into account the possible congestion problem. Based on simulations the analytical form of blocking time for pickers was proposed. The study shows that the advantages of each additional merged order is getting smaller. Using ABC classification for storage location assignment in warehouses with low number of aisles and large pick lists can extend the average order-picking time.

Keywords: order-picking, order-batching, optimization, congestion, blocking

JAKUB WITKOWSKI¹ZASTOSOWANIE METOD BADAŃ OPERACYJNYCH
W OCENIE EFEKTYWNOŚCI MECHANIZMÓW SELEKCJI
W ORGANIZACJACH CHARYTATYWNYCH²

1. WSTĘP

Współczesne społeczeństwa angażują się w działalność dobroczynną i charytatywną przeznaczając na nią znaczne fundusze (np. w samych Stanach Zjednoczonych w 2000 roku wartość pomocy charytatywnej została oszacowana na 2% wartości PKB tego kraju (Wright, 2001)). Dlatego też działalność mająca na celu pomoc osobom potrzebującym jest przedmiotem badań zarówno z obszaru ekonomii jak i socjologii. Badane są między innymi: motywacje osób zaangażowanych w działalność dobroczynną (por. np. Ariely i inni, 2009), przyczyny decyzji o wsparciu danej organizacji charytatywnej (van Iwaarden i inni, 2009), oraz zależności pomiędzy kształtem systemu podatkowego a udzielaniem pomocy charytatywnej (np. Clotfelter, 1980). Ważnym aspektem badań jest też próba określenia do kogo pomoc powinna trafiać, ponieważ bieda może dotyczyć wielu wymiarów życia (np. warunków materialnych takich jak mieszkanie czy dochód, ale także aspektów niematerialnych – np. więzi społecznych), dlatego trudno jednoznacznie ją zdefiniować (por. Haveman, Bershadker, 1998; Perry, 2002). W zależności od przyjętej definicji biedy, programy charytatywne mogą mieć na celu pomoc różnym grupom społecznym zagrożonym biedą, lub uzyskanie zróżnicowanych efektów, np. wsparcie osób posiadających niski dochód, albo przekazanie środków osobom mających największą szansę ekonomicznego usamodzielnienia się.

Wiele badań koncentruje się na mierzeniu skuteczności pomocy charytatywnej. Jednym z możliwych sposobów pomiaru jest ocena, czy pomoc trafiła do właściwych osób (por. Cornia, Stewart, 1993). Zakłada się, że nie każdy potencjalny beneficjent wykorzysta otrzymane zasoby wg zamierzeń darczyńcy, albo że pewna wybrana grupa osób najlepiej skorzysta z dostępnej pomocy (por. Morrissey, 2006; Bougheas i inni 2007; Das i inni, 2005). W związku z tym wiele publikacji naukowych skupia się na analizie doboru grupy docelowej pomocy, tak aby pomoc ta była jak najbardziej efektywna. W takich analizach zakłada się, że nie jest znany dokładny dochód (najczęściej

¹ Szkoła Główna Handlowa w Warszawie, Instytut Ekonometrii, Zakład Wspomagania i Analizy Decyzji, Al. Niepodległości 162, 02-554 Warszawa, Polska, email: jzmwitkowski@gmail.com.

² Pokazane tutaj rezultaty zostały zaprezentowane na XXXV Konferencji Naukowej „Metody i Zastosowania Badań Operacyjnych” im. Profesora Władysława Bukietyńskiego (MZBO'16).

ta zmienna jest wykorzystywana do pomiaru biedy), ale znany jest rozkład dochodu w pewnych grupach populacji (por. Dasgupta, Kanbur, 2005). Na podstawie tej informacji przydzielana jest pomoc w taki sposób, aby zminimalizować zadany wskaźnik mierzący ubóstwo. Inne podejście do problemu prezentuje Bougheas i inni (2007), pokazując problem przydziału środków jako model pryncypała-agenta. W modelu tym, darczyńcy rozdzielają środki pomiędzy agentów o różnych charakterystykach w taki sposób aby zmaksymalizować łączną konsumpcję. Otrzymanie pomocy wymaga od agentów podjęcia działania, które pociąga za sobą koszt. Ponieważ nie każdy typ agenta chce ponosić ten koszt, następuje samoczynna selekcja. W rzeczywistości często możliwa jest też sytuacja, kiedy organizacja charytatywna albo darczyńca, a nie potencjalny beneficjent, ponosi koszt weryfikacji czy dana osoba spełnia kryteria programu pomocy charytatywnej (Dynarski, Scott-Clayton (2006) opisują ten problem w odniesieniu do procesu przyznawania stypendiów dla studentów).

Ważnym elementem procesu przyznawania pomocy charytatywnej jest zdefiniowanie grupy, która ma zostać objęta pomocą oraz weryfikacja czy osoby z tej grupy najlepiej wykorzystają otrzymane środki. Jednak przeprowadzenie weryfikacji niesie ze sobą koszty, pojawia się więc problem właściwego rozdziału dostępnych środków. Budżet może zostać rozdzielony pomiędzy przyznawanie pomocy albo weryfikację osób potrzebujących (przykładowo organizacja może część swojego budżetu przeznaczać na wynajęcie pracowników przeprowadzających weryfikację potrzebujących ale oznacza to zmniejszenie kwoty przeznaczanej bezpośrednio na pomoc). Ponieważ w literaturze poświęconej działalności dobroczynnej temat ten nie jest wystarczająco omówiony, celem tej pracy jest sformułowanie modelu pozwalającego opisać i rozwiązać problem rozdziału środków pomiędzy pomoc i weryfikację.

2. OPIS PROBLEMU

Organizacja charytatywna jest pośrednikiem pomiędzy społeczeństwem a osobami potrzebującymi w tym sensie, że rozdysponowuje otrzymane od społeczeństwa środki pieniężne pomiędzy osoby, które wg swoich kryteriów uzna za potrzebujące. Ponieważ społeczeństwo dotuje takie organizacje, których kryteria przyznawania pomocy uważa za właściwe (tzn. pomagające wypełnić ważne dla społeczeństwa cele takie jak: zmniejszenie nierówności, aktywizacja osób długotrwale bezrobotnych, etc.) to przyznawanie pomocy w taki sposób jest społecznie użyteczne. W przypadku gdyby pomoc została przyznana wbrew tym kryteriom, takie działanie zmniejsza dobrostan społeczeństwa (np. ze względu na koszt utraconych możliwości).

W celu sprawdzenia czy dana osoba spełnia kryteria przyznawania pomocy organizacja charytatywna może podjąć działania mające na celu weryfikację tej osoby (zebranie informacji na temat danej osoby), co jednak jest kosztowne (czas pracowników lub wolontariuszy, który mógłby być przeznaczony na inne czynności). Dlatego organizacja staje przed wyborem czy lepiej z punktu widzenia społeczeństwa przyznać pomoc

mniejszej liczbie zweryfikowanych osób, czy też lepiej zrezygnować z weryfikacji, ale za to przyznać pomoc większej grupie osób ryzykując, że część tej pomocy trafi do osób niespełniających założonych kryteriów.

3. MODEL

Opisany problem decyzyjny można przedstawić za pomocą modelu optymalizacyjnego. W modelu, w którym rozważany jest jeden okres można wyróżnić następujące części:

- potrzebujący: grupa osób którym będzie udzielona pomoc materialna. Załóżmy, że jest N osób, które mogą być brane pod uwagę przy rozdziale środków pomocowych. Można wyróżnić dwa typy potrzebujących: osoby, które spełniają kryteria włączenia do programu pomocy charytatywnej (ich odsetek wynosi n ($n \in (0,1)$)) oraz osoby nie spełniające tych kryteriów (jest ich $1 - n$). Typy poszczególnych osób są niejawne dla społeczeństwa i organizacji charytatywnej.

Średnia wartość pomocy otrzymywana przez osobę potrzebującą wynosi p . Przydzielenie pomocy osobie z pierwszej grupy spowoduje wzrost użyteczności społeczeństwa o $f(p)$, zaś udzielenie jej osobie z drugiej grupy spadek o $kf(p)$ gdzie $k > 1$. Taka asymetria wynika z różnego postrzegania zysku i straty (teoria perspektywy (Kahneman, Tversky, 1979)).

Użyteczność społeczeństwa z przydzielanej powinna rosnąć wraz ze wzrostem średniej kwoty pomocy, jednak krańcowa użyteczność powinna być malejąca (tak aby zapobiec koncentracji przydzielanej pomocy wśród kilku potrzebujących osób). Dlatego o $f(p)$ należy założyć, że spełnia warunki: $f'(p) > 0$ i $f''(p) < 0$.

- organizacja charytatywna: pośrednik, decydujący o tym, jak rozdzielić otrzymany budżet. Organizacja może go wykorzystać na trzy sposoby: przyznawać pomoc poprzedzoną weryfikacją, przyznawać pomoc bez weryfikacji, albo pozostawić otrzymane środki w budżecie. Weryfikacja wiąże się z poniesieniem przez organizację pewnego kosztu, w celu poznania prawdziwego typu danej osoby (czy kwalifikuje się do programu czy nie). Liczbę osób zweryfikowanych oznaczmy jako Z , a liczbę osób niezwyfikowanych jako R . Zakładając pełną skuteczność procesu weryfikacji, spośród zweryfikowanych Z osób średnio nZ otrzyma pomoc. Tak więc użyteczność uzyskana z grupy osób zweryfikowanych wyniesie $nf(p)Z$. Oznaczając średni koszt weryfikacji jednej osoby jako c łączny koszt przydziału pomocy i weryfikacji można zapisać jako $cZ + npZ$.

W drugim przypadku (brak weryfikacji) wszystkie osoby niezwyfikowane otrzymają pomoc, ale nR osób przyniesie użyteczność dodatnią, a $(1 - n)R$ ujemną (czyli łączną użyteczność z przyznania pomocy bez weryfikacji można zapisać jako $nf(p)R - (1 - n)kf(p)R$). Łączny koszt przyznawania pomocy osobom niezwyfikowanym wyniesie pR .

Trzeci rozpatrywany przypadek wiąże się z tym, że istnieje pewien próg opłacalności ekonomicznej pomocy i jeżeli nie zostanie on przekroczony dla społeczeństwa bardziej opłacalne jest zainwestowanie zgromadzonych środków w inny sposób.

Taka inwestycja będzie opłacalna dla społeczeństwa w sytuacji kiedy albo koszty pomocy znacznie przekraczają jej wartość (tzn. stosunek $\frac{c}{p}$ jest wystarczająco wysoki), albo kiedy odsetek osób potrzebujących kwalifikujących się do programu pomocowego jest bardzo niski. Użyteczność związana z pozostawianiem pieniędzy w budżecie jest opisana przez funkcję $g(p, c, n)$, spełniającą następujące warunki:

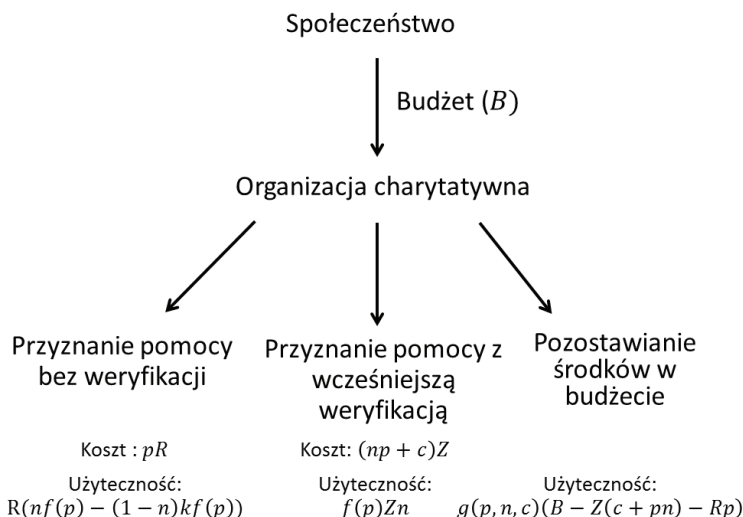
- wartość funkcji rośnie wraz ze wzrostem kosztu weryfikacji ($g'_c > 0$), a spada wraz ze wzrostem n oraz wzrostem p ($g'_p < 0$ oraz $g'_n < 0$),
- dla wartości c/p większych od pewnego progu (d) wartość funkcji $g(p, c, n)$ przyrasta szybciej ze względu na wzrost c niż spada ze względu na wzrost p . Dla wartości p/c niższych od progu d zachodzi zależność odwrotna.

$$\left(\frac{c}{p} \geq d \Rightarrow |g'_c(p, c, n)| - |g'_p(p, c, n)| > 0 \text{ oraz} \right.$$

$$\left. \frac{c}{p} < d \Rightarrow |g'_c(p, c, n)| - |g'_p(p, c, n)| < 0, \text{ gdzie } d \geq 0 \right)$$

Funkcja łącznej użyteczności społeczeństwa, inkorporującej w sobie wszystkie te składniki, będzie oznaczana jako $U(Z, R)$.

Podsumowując, organizacja charytatywna może podjąć trzy decyzje: przyznanie pomocy potrzebującym bez weryfikacji, przyznanie pomocy z wcześniejszą weryfikacją, albo pozostawienia środków w budżecie. Ich koszt oraz wpływ na użyteczność społeczeństwa został schematycznie przedstawiony na rysunku 1.



Rysunek 1. Schemat modelu rozdziału środków w organizacji charytatywnej

Źródło: opracowanie własne.

Z tak przyjętych założeń wynika model matematyczny opisujący podejmowanie decyzji przez organizację charytatywną:

$$U(Z, R) = (Znf(p)) + (Rnf(p) - R(1 - n)kf(p)) + g(p, n, c)(B - Z(c + pn) - Rp) \rightarrow \max,$$

pw.

$$Rp + Z(np + c) \leq B \quad (1)$$

(wydatki na pomoc nie mogą przekroczyć budżetu),

$$0 \leq R \leq \frac{B}{p} \quad (2)$$

(liczba osób, którym udzielono pomocy bez weryfikacji jest ograniczona przez kwotę budżetu),

$$0 \leq Z \leq \frac{B}{c + np} \quad (3)$$

(liczba osób, którym udzielono pomocy z wcześniejszą weryfikacją jest ograniczona przez kwotę budżetu).

Korzystając z twierdzenia Kuhna–Tuckera rozwiązanie zadanego problemu można uzyskać rozwiązując następujący układ równań:

$$\left\{ \begin{array}{l} f(p)(n - (1 - n)k) - g(p, n, c)p + \lambda_1 p + \lambda_3 - \lambda_4 = 0, \\ f(p)(n) - (np + c)g(p, c, n) + \lambda_1(c + np) + \lambda_2 - \lambda_5 = 0, \\ \lambda_1(B - Z(c + np) - Rp) = 0, \\ \lambda_2 Z = 0, \\ \lambda_3 R = 0, \\ \lambda_4\left(\frac{B}{p} - R\right) = 0, \\ \lambda_5\left(\frac{B}{np + c} - Z\right) = 0. \end{array} \right. ,$$

Ze zbioru rozwiązań tego układu równań, tylko trzy z nich spełniają warunki ograniczające modelu opisującego działalność organizacji charytatywnej. Rozwiązania te są następujące:

- S1: cały budżet zostaje przeznaczony na pomoc bez wcześniejszej weryfikacji osób potrzebujących ($R = \frac{B}{p}; Z = 0$),
- S2: cały budżet zostaje przeznaczony na pomoc z wcześniejszą weryfikacją osób potrzebujących ($R = 0; Z = \frac{B}{c+pn}$),
- S3: całość środków pozostaje w budżecie ($R = 0; Z = 0$).

Dla każdej z wartości kombinacji parametrów n , p i c inne rozwiązanie może dawać najwyższą wartość użyteczności, tzn. to które z rozwiązań będzie optymalne, zależy od wartości tych parametrów. Porównując wartość funkcji użyteczności dla poszczególnych rozwiązań można wyprowadzić warunki pozwalające na stwierdzenie, które z rozwiązań modelu będzie optymalne:

- jeżeli spełniony jest układ nierówności:

$$\left\{ \begin{array}{l} n - k(1 - n) \geq \frac{np}{c + np}, \\ \frac{f(p)}{p} (n - k(1 - n)) \geq g(p, n, c), \end{array} \right.$$

to najbardziej efektywnie kosztowo jest przydzielanie pomocy bez weryfikacji, a więc optymalne jest rozwiązanie S1;

- jeżeli spełnione są nierówności:

$$\left\{ \begin{array}{l} \frac{np}{c + np} \geq n - k(1 - n), \\ \frac{nf(p)}{c + np} \geq g(p, n, c), \end{array} \right.$$

to rozwiązaniem optymalnym jest S2, czyli przydział środków po wcześniejszej weryfikacji;

- optymalną decyzją jest pozostawienie środków w budżecie (rozwiązanie S3), jeżeli spełniony jest następujący układ równań:

$$\left\{ \begin{array}{l} g(p, n, c) \geq \frac{f(p)}{p} (n - k(1 - n)), \\ g(p, n, c) \geq \frac{nf(p)}{c + np}. \end{array} \right.$$

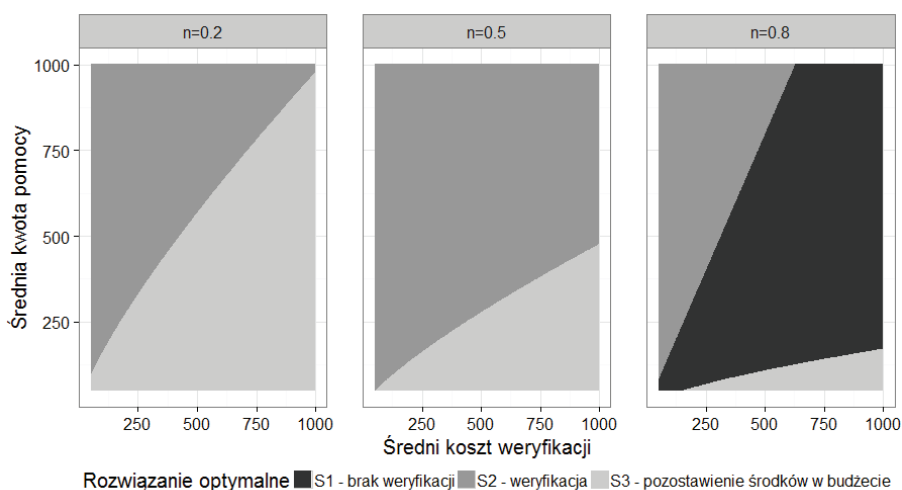
Dla lepszego zobrazowania zależności pomiędzy wartościami parametrów, a optymalnością poszczególnych rozwiązań na rysunku 2 przedstawiono, które rozwiązanie

jest optymalne przy zadanych wartościach parametrów c , p i n . Na potrzeby dalszego wywodu przyjęte zostały postacie funkcyjne dla $f(p)$ i $g(p, c, n)$ spełniające opisane wcześniej warunki. Niech $f(p) = \sqrt{p}$ oraz $g(p, c, n) = \frac{c}{np^2}$.

W przypadku, kiedy odsetek osób kwalifikujących się do programu jest niski (0,2) optymalne może być rozwiązanie S2 albo S3, co jest przedstawione na lewym panelu. Zgodnie z intuicją, o tym, które z rozwiązań będzie optymalne, decyduje relacja pomiędzy kwotą pomocy a kosztem weryfikacji.

Wraz ze wzrostem odsetka osób kwalifikujących się do programu 0,5 (środkowy panel) sytuacja nie zmienia się – dalej rozwiązaniem optymalnym może być S2 lub S3, ale tym razem inna jest proporcja kwoty pomocy i kosztu weryfikacji, przy której będą wybierane poszczególne rozwiązania (na korzyść S2).

W przypadku, kiedy odsetek osób kwalifikujących się do programu jest wysoki i wynosi 0,8 (prawy panel) rozwiązaniem optymalnym może okazać się także S1. Rozwiązania S2 i S3 są optymalne tylko dla niskich lub wysokich wartości stosunku parametrów p i c , zaś w pozostałych przypadkach S1 staje się rozwiązaniem optymalnym.



Rysunek 2. Optymalność rozwiązania w zależności od wartości parametrów (dla $k = 1,2$)

Źródło: opracowanie własne.

Podsumowując, z modelu wynika, że weryfikacja osób zgłaszających się do programu pomocy charytatywnej jest opłacalna z punktu widzenia społeczeństwa jeżeli koszt weryfikacji nie jest zbyt wysoki w stosunku do kwoty przydzielanej pomocy. Na to, czy w danej sytuacji weryfikacja jest opłacalna ma też wpływ odsetek osób kwalifikujących się do programu – przy niskim odsetku osób kwalifikujących się do programu stosunek kosztu weryfikacji do kwoty pomocy musi być niższy aby

weryfikacja była opłacalna, niż w sytuacji w której odsetek osób kwalifikujących się do programu jest wysoki. Także, przy wysokich wartościach opłacalne staje się przydzielanie pomocy bez wcześniejszej weryfikacji.

4. ANALIZA WRAŻLIWOŚCI

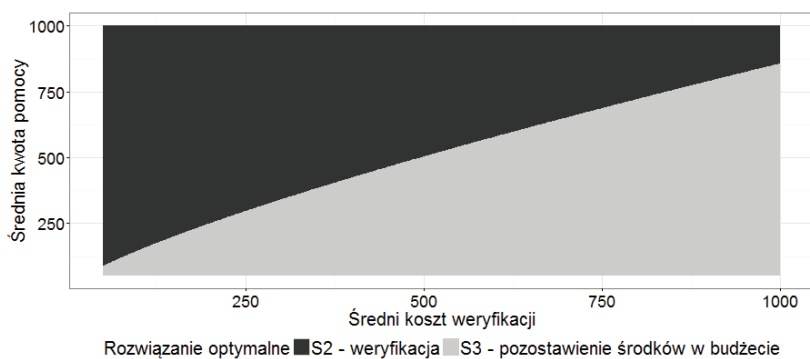
Dotychczas uzyskane wyniki są prawdziwe przy założeniu, że w każdej sytuacji znane są wartości parametrów n , p i c . Tak jest na pewno w sytuacji, kiedy program pomocowy został przeprowadzony i kiedy można dokładnie oszacować te parametry na podstawie danych empirycznych. Jednak dużo ważniejszym jest pytanie czy można ocenić opłacalność weryfikacji osób potrzebujących przed przeprowadzeniem danego programu? W przypadku parametrów opisujących koszt weryfikacji lub kwotę pomocy są one albo znane z góry albo mogą być w łatwy sposób oszacowane. Jednak w przypadku parametru n , trudno jest przewidzieć jaka będzie jego dokładna wartość. Na podstawie danych zebranych przy realizacji podobnych programów można podać przybliżoną wartość n , ale rzeczywista wartość tego parametru jest trudna do oszacowania na tej podstawie.

Do podejmowania decyzji w sytuacji, kiedy jeden z parametrów nie jest znany, a jedynie można przyjąć pewne założenia co do wartości, które może on przyjąć (czyli zakłada się, że pochodzi on z pewnego rozkładu) wykorzystuje się podejście bayesowskie (por. Berger, 1985). Przyjmując, że organizacja podejmuje jedną z trzech akcji (tzn. wdraża jedno z trzech rozwiązań: S1, S2 lub S3) dla każdego zestawu parametrów p i c istnieje prawdopodobieństwo $\pi(S_i)$ oznaczające że i -te rozwiązanie jest optymalne, zależne od stanu świata, czyli wartości n . Także dla każdego zestawu parametrów p i c można określić wartość funkcji straty $L(S_i, n) = |U(S_{opt}|n) - U(S_i|n)|$ w przypadku nie wybrania rozwiązania optymalnego. Każdej decyzji można przyporządkować wartość oczekiwanej straty z podjęcia danej decyzji: $\int L(S_i, n)\pi(n)dn$. W takich warunkach optymalną decyzją będzie ta, która wiąże się z najmniejszą oczekiwaną stratą.

Jeżeli przyjąć, że każda wartość n jest tak samo prawdopodobna (tzn. nie ma informacji, które pozwalałyby przyjąć założenia o możliwym dokładniejszym zakresie zmienności parametru n) to można założyć, że $n \sim U(0,1)$. Rozkład rozwiązań optymalnych (tzn. przynoszących najmniejszą oczekiwaną stratę) przedstawia rysunek 3.

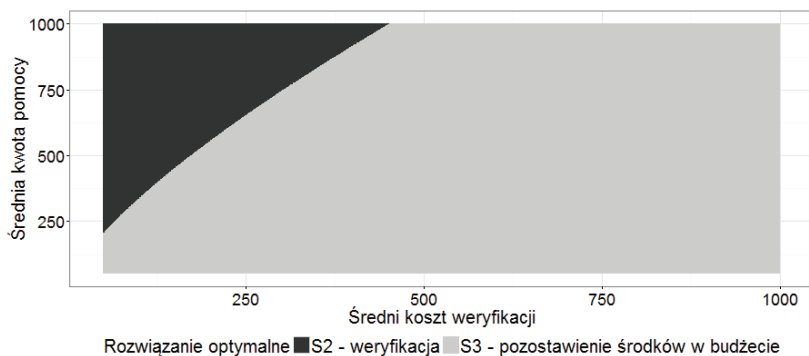
Na podstawie rysunku 3 można zauważyć, że przy niskiej wartości ilorazu kwoty pomocy i kosztu weryfikacji optymalne jest rozwiązanie S3, a przy wysokiej – S2. Uzyskane wyniki są bardzo podobne do tych, które zostały uzyskane przy założeniu o niskiej wartości parametru n (por. rysunek 2).

Otrzymany wynik jest zależny od przyjętej funkcji straty. Gdyby zamiast absolutnej funkcji straty przyjąć kwadratową $L(S_i, n) = (U(S_{opt}|n) - U(S_i|n))^2$, to rozwiązanie S2 byłoby wybierane dużo rzadziej, co przedstawia rysunek 4.



Rysunek 3. Optymalność rozwiązania w zależności od wartości parametrów (dla $k = 1,2$), przy założeniu $n \sim U(0,1)$

Źródło: opracowanie własne.



Rysunek 4. Optymalność rozwiązania w zależności od wartości parametrów (dla $k = 1,2$), przy założeniu $n \sim U(0,1)$ przy założeniu kwadratowej funkcji straty

Źródło: opracowanie własne.

Podsumowując, brak wiedzy o dokładnej wartości parametru n eliminuje z rozważanych rozwiązań możliwość przydzielania pomocy bez wcześniejszej weryfikacji. Rozważane mogą być rozwiązania S2 lub S3, ale to które z nich będzie rozwiązaniem optymalnym przy konkretnych wartościach parametrów p i c oraz założeniu $n \sim U(0,1)$ zależy od przyjętej funkcji straty. Przy wykorzystaniu funkcji kwadratowej częściej optymalne będzie rozwiązanie S3 niż S2. Weryfikacja będzie rozwiązaniem optymalnym tylko dla sytuacji, w których kwota pomocy jest dużo wyższa od kosztu weryfikacji. W przypadku wykorzystania jako funkcji straty wartości absolutnej, proporcja rozwiązań zmieni się – S2 będzie wybierane częściej (wartość stosunku p i c przy którym S2 będzie wybierane jako optymalne zmniejszy się w porównaniu do wartości tego stosunku przy kwadratowej funkcji straty).

5. STUDIUM PRZYPADKU

Jednym z projektów charytatywnych w Polsce, w którym silny nacisk jest kładziony na weryfikację osób zgłaszanych jako potrzebujące jest Szlachetna Paczka. W tym projekcie, każda zgłoszona rodzina jest weryfikowana przez wolontariuszy pod względem wewnętrznych kryteriów projektu (dotyczących m.in. postawy życiowej). W 2015 roku wolontariusze odwiedzili w całej Polsce prawie 32 tysiące rodzin. Z tej liczby 19580 rodzin zostało włączonych do projektu i otrzymało pomoc, czyli przygotowywane przez darczyńców paczki, mające zaspokajać indywidualne potrzeby rodzin. Średnia wartość paczki w 2015 roku wyniosła 2583 PLN (por. Stowarzyszenie WIOSNA, 2015).

Dane z 2015 roku pozwalają wyznaczyć wartości parametrów $n \approx 0.63$ oraz $p \approx 2583$. Trudniej jest wyznaczyć dokładnie wartość średniego kosztu weryfikacji, ponieważ jest ona przeprowadzana przez wolontariuszy. Zakładając, że wolontariusz musi poświęcić na weryfikację i przeprowadzenie przez projekt jednej rodziny średnio od 4 do 10 godzin, przy średnim wynagrodzeniu brutto w 2015 na poziomie 3899 PLN (por. GUS, 2016) koszt weryfikacji rodziny wyniósłby od około 100 do 250 PLN (pomijając koszt stworzenia oraz utrzymania systemu informatycznego, szkolenia wolontariuszy oraz innych środków wspomagających proces weryfikacji rodzin). W takiej sytuacji okazuje się, że weryfikacja rodzin jest ekonomicznie uzasadniona – dla takich wartości parametrów model wskazuje, że społeczeństwo uzyska najwyższą użyteczność. Ten wniosek jest prawdziwy dla bardzo szerokiego przedziału c (od 0 do 7121,6 PLN), tak więc nawet ewentualny błąd w estymacji kosztu nie ma wpływu na otrzymane wyniki.

6. PODSUMOWANIE

W tym artykule został przedstawiony model opisujący problem rozdziału środków przez organizację charytatywną. Pokazuje on, jaki sposób podziału środków (przydzielanie pomocy wszystkim zgłoszonym osobom, przydzielenie pomocy poprzedzone weryfikacją albo pozostawianie środków w budżecie) jest najlepszy dla społeczeństwa w zależności od parametrów programu charytatywnego: średniej kwoty przydzielanej pomocy, kosztu weryfikacji oraz odsetka osób potrzebujących spełniających założenia programu.

Uzyskane wyniki z modelu pokazują, że decyzja o alokacji środków powinna być zależna od proporcji wartości pomocy do kosztów weryfikacji. Graniczna wartość tej proporcji jest uzależniona od odsetka osób potrzebujących spełniających kryteria programu. W sytuacjach kiedy stosunek kwoty pomocy do kosztu weryfikacji jest wysoki, najlepszą dla społeczeństwa decyzją będzie przyznawanie pomocy z wcześniejszą weryfikacją. W odwrotnym przypadku (tzn. kiedy stosunek kwoty pomocy do kosztu weryfikacji jest niski) opłacalne jest zatrzymanie pieniędzy w budżecie. Przyznawanie

pomocy bez weryfikacji okazuje się rozwiązaniem optymalnym tylko w sytuacji kiedy odsetek osób spełniających kryteria programu jest wysoki.

Przeprowadzona analiza wrażliwości modelu, polegająca na uzmiennieniu parametru opisującego odsetek osób spełniających kryteria programu pokazuje, że w sytuacji kiedy nie jest znana dokładana wartość tego parametru rozwiązaniami optymalnymi mogą być albo przydzielanie pomocy z wcześniejszą weryfikacją albo zatrzymanie środków pomocowych w budżecie organizacji. Kryterium wyboru pomiędzy tymi rozwiązaniami jest stosunek kwoty pomocy do kosztu weryfikacji (a jego wartości graniczne zależą od postaci funkcji straty przyjętej w analizie wrażliwości).

Przedstawiony model pokazuje, że decyzja o tym czy przydzielanie pomocy charytatywnej powinno być poprzedzone wcześniejszą weryfikacją osób potrzebujących zależy od warunków danego programu charytatywnego. Okazuje się, że w pewnych przypadkach najlepszym dla społeczeństwa rozwiązaniem jest nieprzydzielenie pomocy ze względu na jej nieopłacalność ekonomiczną.

LITERATURA

- Ariely D., Bracha A., Meier S., (2009), Doing Good or Doing Well? Image Motivation and Monetary Incentives in Behaving Prosocially, *The American Economic Review*, 99 (1), 544–555.
- Berger J. O., (1985), Statistical Decision Theory and Bayesian Analysis, *Springer Science & Business Media*.
- Bougheas S., Dasgupta I., Morrissey O., (2007), Tough Love or Unconditional Charity?, *Oxford Economic Papers*, 59 (4), 561–582.
- Clotfelter C. T., (1980), Tax Incentives and Charitable Giving: Evidence From a Panel of Taxpayers, *Journal of Public Economics*, 13 (3), 319–340.
- Cornia G. A., Stewart F., (1993), Two Errors of Targeting, *Journal of International Development*, 5 (5), 459–496.
- Das J., Do Q. T., Özler B., (2005), Reassessing Conditional Cash Transfer Programs, *The World Bank Research Observer*, 20 (1), 57–80.
- Dasgupta I., Kanbur R., (2005), Community and Anti-Poverty Targeting, *The Journal of Economic Inequality*, 3 (3), 281–302.
- Dynarski S. M., Scott-Clayton J. E., (2006), The Cost of Complexity in Federal Student Aid: Lessons from Optimal Tax Theory and Behavioral Economics, *National Bureau of Economic Research*, 59, 319–356.
- GUS (2016), Zatrudnienie i wynagrodzenia w gospodarce narodowej w 2015 roku, <http://stat.gov.pl/obszary-tematyczne/rynek-pracy/pracujacy-zatrudnieni-wynagrodzenia-koszty-pracy/zatrudnienie-i-wynagrodzenia-w-gospodarce-narodowej-w-2015-roku,1,21.html>, data dostępu: 14/11/2016.
- Haveman R. H., Bershadker A., (1998), The “Inability to Be Self-Reliant” as an Indicator of Poverty: Trends in the United States, 1975-1995, *University of Wisconsin--Madison, Institute for Research on Poverty*, 47, 335–60.
- Kahneman D., Tversky A., (1979), Prospect Theory: An Analysis of Decision Under Risk, *Econometrica, Journal of the Econometric Society*, 47, 263–291.
- Morrissey O., (2006), Fungibility, Prior Actions, and Eligibility for Budget Support, *Budget Support as More Effective Aid*, 333–350.
- Perry B., (2002), The Mismatch Between Income Measures and Direct Outcome Measures of Poverty, *Social Policy Journal of New Zealand*, 19, 101–127.
- Stowarzyszenie WIOSNA (2015), Raport o polskiej biedzie, http://www.szlachetnapaczka.pl/upload/rob2015/raport_o_biedzie.pdf, data dostępu: 14/11/2016.

- van Iwaarden J., van der Wiele T., Williams R., & Moxham C., (2009), Charities: How Important is Performance to Donors?, *International Journal of Quality & Reliability Management*, 26 (1), 5–22.
- Wright K., (2001), Generosity vs. Altruism: Philanthropy and Charity in the United States and United Kingdom, *Voluntas: International Journal of Voluntary and Nonprofit Organizations*, 12 (4), 399–416.

ZASTOSOWANIE METOD BADAŃ OPERACYJNYCH W OCENIE EFEKTYWNOŚCI MECHANIZMÓW SELEKCJI W ORGANIZACJACH CHARYTATYWNYCH

Streszczenie

W artykule badana jest ekonomiczna efektywność działań mających na celu weryfikację potencjalnych beneficjentów programów charytatywnych. Taka weryfikacja ma na celu ograniczenie liczby osób otrzymujących pomoc charytatywną, które nie spełniają kryteriów jej przyznania. Jednocześnie, wiąże się ona z kosztami ponoszonymi przez organizatora programu. W artykule zaproponowano model teoretyczny opisujący to zjawisko. Do przeprowadzania analizy modelu wykorzystano metody optymalizacyjne oraz symulacyjne. Wynika z niej, że w zależności od parametrów programu charytatywnego optymalna może być jedna z trzech decyzji: przyznanie pomocy bez wprowadzania weryfikacji, przyznawanie pomocy po uprzedniej weryfikacji lub przeznaczenie środków pomocowych na inny cel. Optymalność rozwiązania zależy od charakterystyk programu, takich jak średnia kwota pomocy, średni koszt weryfikacji oraz odsetek osób potrzebujących spełniających założenia programu. W zależności od tych zmiennych organizacja charytatywna powinna podjąć decyzję o ewentualnym wprowadzeniu weryfikacji.

Słowa kluczowe: pomoc charytatywna, optymalizacja, programowanie liniowe

ASSESSING SELECTION MECHANISMS IN CHARITY ORGANIZATIONS USING OPERATIONAL RESEARCH METHODS

Abstract

The aim of this paper is to assess efficiency of verification policies used in charity organizations. Verification policy can be introduced to ensure that aid is granted only to people who are eligible for it. However, it bears costs for the program's organizers. The paper provides a theoretical model assessing economic efficiency of verification policy using optimization and simulation methods. Depending on characteristics of the program one of three decisions might be optimal: granting aid without verification, granting aid only to positively verified people or not granting aid at all. Therefore, in order to make a decision about introducing verification policy charity organization should analyze parameters of the program such as: average amount of aid, average cost of verification and percentage of people eligible to obtain help in the program.

Keywords: charity aid, optimization, linear programming