

AKADEMIA EKONOMICZNA IM. KAROLA ADAMIECKIEGO

Prace naukowe

MODELOWANIE PREFERENCJI A RYZYKO '10

**Praca zbiorowa pod redakcją naukową
Tadeusza Trzaskalika**



KATOWICE 2010

Komitet Redakcyjny

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

Recenzenci

**Ignacy Kaliszewski
Donata Kopańska-Bródka
Dorota Kuchta
Honorata Sosnowska
Józef Stawicki
Włodzimierz Szkutnik**

Redaktor

Jadwiga Popławska-Mszyca

Projekt okładki

Urszula Grendys

Niniejsza publikacja została pierwotnie wydana w wersji drukowanej/komercyjnej.
Publikacja została ponownie udostępniona w modelu otwartego dostępu (Open Access)
i jest dostępna bezpłatnie na licencji Creative Commons Uznanie autorstwa 4.0
Międzynarodowa (CC BY 4.0), <https://creativecommons.org/licenses/by/4.0/legalcode.pl>

© Copyright by Wydawnictwo Akademii Ekonomicznej w Katowicach 2010
(wersja drukowana)

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2025
(wersja elektroniczna)



ISBN 978-83-7246-566-5 (wersja drukowana)
ISBN 978-83-7875-958-4 (wersja elektroniczna)

**WYDAWNICTWO AKADEMII EKONOMICZNEJ
IM. KAROLA ADAMIECKIEGO W KATOWICACH**

ul. 1 Maja 50, 40-287 Katowice, tel. 032 257-76-35, fax 032 257-76-43
www.ae.katowice.pl, e-mail: wydawnictwo@ae.katowice.pl

SPIS TREŚCI

WPROWADZENIE	7
OPTYMALNE DECYZJE	17
Paweł Baran: ZASTOSOWANIE MODELI PROGRAMOWANIA DYNAMICZNEGO W WARUNKACH ROZMYTYCH W ZARZĄDZANIU MAJĄTKIEM	19
Krzysztof Dmytrów: ZASTOSOWANIE MATEMATYCZNEGO MODELU ZAPASÓW DLA PRODUKTÓW PSUJĄCYCH SIĘ DO GOSPODAROWANIA ZASOBAMI GOTÓWKI	29
Monika Dyduch: WSPÓŁCZYNNIKI TRANSFORMATY FALKOWEJ JAKO NARZĘDZIE GENERUJĄCE PROGNOZĘ PRZEDZIAŁOWĄ SZEREGÓW CZASOWYCH	35
Dorota Gawrońska: SZACOWANIE EFEKTYWNOŚCI PRZEDSIĘ- WZIĘCIA INWESTYCYJNEGO PRZEDSIĘBIORSTWA NA PODSTAWIE WSKAŹNIKA NPVR O ROZMYTYCH PARAMETRACH	47
Paweł Hanczar: TARYFY STREFOWE W PROBLEMACH WYZNACZANIA TRAS POJAZDÓW	67
Katarzyna Jakowska-Suwalska: ZAGADNIENIE LOKALIZACJI OBIEKTÓW MODELOWANYCH JAKO SYSTEMY G/G/1 Z RÓŻNYMI KLASAMI KLIENTÓW	77
Marta Kapica, Cezary Dominiak: OPTYMALIZACJA SYSTEMU MASOWEJ OBSŁUGI NA PRZYKŁADZIE ŚLĄSKIEGO WESOŁEGO MIASTECZKA	87
Jerzy Michnik: ANALIZA KRYTERIÓW WYSTĘPUJĄCYCH W PROCESIE DECYZYJNYM W ZARZĄDZANIU INNOWACJAMI	107
Bogusław Nowak: PLANOWANIE INWESTYCJI BUDOWLANEJ – ANALIZA SYMULACYJNA	119
Sebastian Sitarz: OBJĘTOŚCIOWA ANALIZA WRAŻLIWOŚCI W PROGRAMOWANIU LINIOWYM	133

Adam Sojda: WYKORZYSTANIE ALGORYTMU EWOLUCYJNEGO W ZAGADNIENIU LOKALIZACJI SYSTEMÓW MASOWEJ OBSŁUGI G/G/1	145
Joanna Utkin: KONSEKWENCJE AGREGACJI W SKOŃCZONYM MODELU TERMINOWEJ STRUKTURY STÓP PROCENTO- WYCH	155
Sławomir Żułtak: DYNAMICZNA OPTIMALIZACJA STRATEGII INWESTYCYJNEJ W MODELU HO-LEE	173
ANALIZA PROJEKTÓW	183
Paweł Błaszczuk, Maria B. Kania, Tomasz Błaszczuk: DWUKRYTERIALNA ANALIZA PROJEKTU MODELOWANEGO Z UWZGLĘDNIENIEM BUFORÓW CZASOWYCH I KOSZTOWYCH	185
Monika Fedyczak, Dorota Kuchta: ZRÓWNOWAŻONE PODEJŚCIE W ZARZĄDZANIU RYZYKIEM PROJEKTÓW EUROPEJSKICH	197
Dorota Kuchta: NOWA KONCEPCJA ODPORNOŚCI PLANÓW PROJEKTÓW	209
Dorota Kuchta, Ewa Ptaszyńska: ZARZĄDZANIE RYZYKIEM POPRZECZ UCZENIE SIĘ W PROJEKTACH EUROPEJSKICH	221
Krzysztof S. Targiel: WYKORZYSTANIE OPCJI REALNYCH W SZACOWANIU RYZYKA PRZEKROCZENIA PRZEZ KOSZTY PROJEKTU USTALONEJ WARTOŚCI	231
Wojciech Walczak: ZARZĄDZANIE RYZYKIEM W ZWINNYCH METODYKACH ZARZĄDZANIA PROJEKTAMI	241
TEORIA GIER I NEGOCJACJI	257
Hanna Bury, Dariusz Wagner: WPŁYW WYBORU METODY WYZNACZANIA OCENY GRUPOWEJ NA WYNIK EKSPERTYZY	259
Robert Golański: ZNACZENIE SENATU W POLSKIM SYSTEMIE LEGISLACYJNYM – PRZESTRZENNY MODEL TEORIOGROWY	277
Leszek Klukowski: OGRANICZANIE RYZYKA KOSZTÓW OBSŁUGI DŁUGU PUBLICZNEGO PRZY WYKORZYSTA- NIU ROZWIĄZANIA MINIMAKSOWEGO GRY STRATEGICZNEJ	291

Honorata Sosnowska: REGUŁY FORMALNE RYZYKA PORZĄDKO- WEGO W PRZYPADKU ZDARZEŃ POZYTYWNYCH I NEGATYWNYCH	301
Maciej Wolny: QUASI-KLASYCZNE METODY WSPOMAGANIA RACJONALNEGO PODZIAŁU WYGRANEJ W KOOPERACYJNYCH GRACH Z KOALICJAMI	315
Irena Woroniecka-Leciejewicz: RÓWNOWAGA W GRZE FISKALNO- -MONETARNEJ A PRIORYTETY BANKU CENTRALNEGO I RZĄDU	327

WPROWADZENIE

OPTYMALNE DECYZJE

We wcześniejszych pracach P. Barana zaproponowano wykorzystanie modeli programowania dynamicznego w warunkach rozmytych we wspomaganiu planowania inwestycyjnego – do wyboru oraz szeregowania projektów inwestycyjnych. W pracy ***Zastosowanie modeli programowania dynamicznego w warunkach rozmytych w zarządzaniu majątkiem (P. Baran)*** rozważane jest szersze wykorzystanie tego typu modelowania w zarządzaniu majątkiem firmy lub jednostki samorządowej – rozpatrywane są nowe grupy celów, a oprócz realizacji projektów inwestycyjnych jako poszczególne zadania występują także operacje zbycia elementów majątku.

Gospodarowanie gotówką pod wieloma względami przypomina gospodarowanie zapasami materiałowymi w przedsiębiorstwie. Podobnie jak w przypadku zapasów materiałowych, mamy do czynienia z nabyciem gotówki, jej przechowywaniem, niedoborami oraz psuciem się. Przez nabycie gotówki rozumiemy dostarczenie gotówki do kas w banku bądź do bankomatów, co generuje określone koszty zamawiania. Przechowując gotówkę ponosimy koszty zamrożenia środków, które można utożsamiać z kosztami niewykorzystanych możliwości, które można by osiągnąć poprzez zainwestowanie ich. Koszty niedoboru są najtrudniejsze do oszacowania. Najczęściej można je utożsamiać z utratą wizerunku przedsiębiorstwa bądź utratą klientów. Przez psucie się gotówki rozumiemy utratę jej wartości w wyniku inflacji. W pracy ***Zastosowanie matematycznego modelu zapasów dla produktów psujących się do gospodarowania zasobami gotówki (K. Dmytrów)*** podjęto próbę wykorzystania modelu zapasów dla produktów psujących się do gospodarowania zasobami gotówki w bankomacie banku komercyjnego.

W pracy ***Współczynniki transformaty falkowej jako narzędzie generujące prognozę przedziałową szeregów czasowych (M. Dyduch)*** przeprowadzono prognozę szeregu prezentującego kurs wymiany euro. W tym celu zaproponowano model integrujący sieci neuronowe i transformatę falkową. Dyskretną transformatę falkową przeprowadzono algorytmem „a torus”. Prognozowane wartości kursu euro otrzymano w ostateczności opierając się na

odwrotnej transformaty falkowej dla 8-elemetowych przedziałów czasowych. Przy czym wartości prognozowane generowane są w ostatnim etapie na podstawie współczynników transformaty falkowej wygenerowanych wcześniej przez sieć neuronową.

Przed wyborem przedsięwzięcia inwestycyjnego do realizacji należy dokonać analizy finansowej. Ze względu na niepewność informacyjną dotyczącą wartości przepływów pieniężnych czy stopy dyskontowej, zwłaszcza przy założeniu długiego horyzontu realizacji inwestycji, powinno przyjąć się zmienne uwzględniające tę niepewność. Taką możliwość daje opracowana przez amerykańskiego profesora Lofti Zadeha teoria zbiorów rozmytych, umożliwiającą modelowanie niepewnej informacji. Celem pracy ***Szacowanie efektywności przedsięwzięcia inwestycyjnego przedsiębiorstwa na podstawie wskaźnika NPVR o rozmytych parametrach*** (D. Gawrońska) jest określenie wartości wskaźnika NPVR, bazując na rozmytych ocenach przepływów finansowych oraz stopy dyskontowej, uzyskując rozmytą ocenę wskaźnika NPV oraz PVI. W dalszej kolejności na podstawie procesu defuzyfikacji określa się wartość rzeczywistą wskaźnika finansowego NPVR. Poprzez porównanie wartości wskaźnika dla kolejnych przedsięwzięć inwestycyjnych można określić przedsięwzięcie optymalne.

W pracy ***Taryfy strefowe w problemach wyznaczania tras pojazdów*** (P. Hanczar) rozważono możliwość zastosowania modeli optymalizacyjnych do wyznaczania planu dostaw w przypadku strefowego systemu rozliczania kosztów. W pierwszej części opracowania zaprezentowano wady i zalety strefowego systemu rozliczania oraz jego wykorzystanie w sieciach logistycznych. Następnie przedyskutowano możliwości wspomagania procesu planowania tras w takim systemie z wykorzystaniem metod optymalizacyjnych. Rozważone zostały dwa podejścia, pierwsze wykorzystujące model przydziału oraz drugie korzystające ze sformułowania modelu wyznaczania tras dostaw. Wprowadzone modele zostały użyte do rozwiązania zadania rzeczywistego. Opracowanie kończy opis integracji opracowanej metody z systemem informatycznym wykorzystywanym w badanym przedsiębiorstwie.

W pracy ***Zagadnienie lokalizacji obiektów modelowanych jako systemy G/G/1 z różnymi klasami klientów*** (K. Jakowska-Suwalska) zajęto się zagadnieniem lokalizacji p obiektów traktowanych jako systemy masowej obsługi G/G/1 do obsługi n obiektów umieszczonych w znanej sieci. Każdy obsługiwany obiekt traktowany jest jako odrębne źródło zgłoszeń. W zagadnieniu lokalizacji rozróżnia się więc różne klasy klientów obsługiwanych przez sys-

temy. Zagadnienie rozwiązano metodą podziału i ograniczeń, gdzie za funkcję celu przyjęto minimalizację średniego całkowitego czasu obsługi w systemach. Do obliczeń wykorzystano aproksymację dyfuzyjną.

Współczesne gospodarki charakteryzują się wysokim i stale rosnącym udziałem sektora usług. Istotnym czynnikiem determinującym efektywność w tym sektorze gospodarki jest odpowiednia przepustowość kanałów obsługi klienta, która bezpośrednio wpływa na wielkość przychodów ze sprzedaży. Wykorzystanie teorii systemów masowej obsługi umożliwia uzyskanie odpowiedniego poziomu obsługi klienta przy zadanych nakładach. W pracy *Optymalizacja systemu masowej obsługi na przykładzie Śląskiego Wesołego Miasteczka* (M. Kapica, C. Dominiak) przedstawiono wykorzystanie tej teorii do optymalizacji ilości stanowisk kasowych w badanym obiekcie. Następnie przedstawiony został linowy model optymalizujący funkcjonowanie systemu przy zadanych parametrach.

Zarządzanie innowacjami charakteryzuje się znaczną ilością różnorodnych kryteriów o skomplikowanej strukturze zależności. Wprowadzanie nowych produktów na rynek obarczone jest również wysokim stopniem niepewności, ponieważ decyzje często dotyczą poważnych zmian w technologii lub organizacji, co wiąże się ze znacznymi nakładami inwestycyjnymi. Niepewne są również efekty zmian organizacyjnych i reakcja rynku na innowację. W pracy *Analiza kryteriów występujących w procesie decyzyjnym w zarządzaniu innowacjami* (J. Michnik) przedstawiono charakterystykę kryteriów, próbę ich klasyfikacji oraz wzajemne zależności. Jest to niezbędny etap wstępny, którego wyniki powinny pomóc w wyborze odpowiednich metod wielokryterialnych wspomagających podejmowanie decyzji w zarządzaniu innowacjami.

Współczesne przedsiębiorstwa w miejsce tradycyjnych stylów zarządzania coraz częściej sięgają po metody powszechnie wykorzystywane w zarządzaniu projektami. Do najczęściej stosowanych tradycyjnych metod zarządzania projektami zaliczamy techniki sieciowe (CPM, PERT, GERT, MPM) oraz kompleksowe metodyki uwzględniające aspekty ilościowe i jakościowe (PRINCE2, PMI). W pracy *Planowanie inwestycji budowlanej – analiza symulacyjna* (B. Nowak) zaproponowano wykorzystanie drzewa decyzyjnego do wyboru projektu związanego z inwestycją budowlaną. Dla rozważanych strategii decyzyjnych przeprowadzono eksperymenty symulacyjne, których celem było uzyskanie informacji na temat możliwego dochodu z realizowanego przedsięwzięcia. Proponowany sposób podejścia zilustrowano przykładem projektu branży transportowej opartym na danych rzeczywistych.

Praca *Objętościowa analiza wrażliwości w programowaniu liniowym* (S. Sitarz) omawia problematykę analizy wrażliwości opartą na podejściu objętościowym (volume-based sensitivity analysis). Przedstawione są metody analizowania wrażliwości współczynników funkcji celu w programowaniu liniowym. Podane jest porównanie prezentowanego podejścia z klasycznym modelem analizy wrażliwości, w którym znajdują się przedziały wartości współczynników funkcji celu.

W pracy *Wykorzystanie algorytmu ewolucyjnego w zagadnieniu lokalizacji systemów masowej obsługi G/G/1* (A. Sojda) przedstawiono algorytm ewolucyjny wykorzystany do znalezienia optymalnej lokalizacji p systemów masowej obsługi G/G/1 obsługujących n obiektów będących odrębnymi źródłami zgłoszeń. Algorytm ewolucyjny wykorzystuje algorytm ewolucji różnicowej wzbogacony heurystykami przeszukiwań lokalnych w celu poprawy znalezionej już rozwiązania. Zastosowanie heurystyk poprawia efektywność działania algorytmu pozwalając uzyskać porównywalne rozwiązania w krótszym czasie.

W pracy *Konsekwencje agregacji w skończonym modelu terminowej struktury stóp procentowych* (J. Utkin) Skończony dwumianowy model terminowej struktury stóp procentowych jest poddany agregacji metodą Klaassena. Najpierw zdefiniowano agregację obligacji zerokuponowych, wyodrębniając szczególną postać prawdopodobieństwa przejścia i czynnika dyskontującego w modelu zagregowanym. Udowodniony został związek między nierównościami Jarrova w modelu niezagregowanym i zagregowanym. Dotyczy on w szczególności dwóch podstawowych modeli terminowej struktury: Pedersena-Shiu-Thorlaciusa i Sandmanna-Sondermanna. Zaproponowana została pewna zasada ustalania subiektywnego prawdopodobieństwa w modelu zagregowanym. Następnie omówiono wyniki rozszerzenia agregacji na dowolne zobowiązanie dane na końcu okresu agregacji.

Celem pracy *Dynamiczna optymalizacja strategii inwestycyjnej w modelu Ho-Lee* (S. Żułtak) jest przedstawienie analitycznych rozwiązań problemu poszukiwania optymalnych, dynamicznych strategii inwestycyjnych w dyskretnym i skończonym modelu stóp procentowych. Optymalna, dynamiczna strategia inwestycyjna definiowana jest jako sekwencja jednookresowych i warunkowanych aktualnym stanem modelu decyzji alokacji majątku pomiędzy obarczone ryzykiem aktywa inwestycyjne w sposób minimalizujący funkcję straty, będącą miarą niedotrzymania losowego zobowiązania. Modelem

ewolucji stóp procentowych rozważanym w pracy jest specyfikacja zaproponowana przez Ho-Lee, natomiast minimalizowana funkcja straty zapożyczona została od Cvitanica i Karatzasa. Opracowanie stanowi kontynuację i rozszerzenie na przypadki dynamiczne badań przedstawionych w innych pracach autora poświęconych zabezpieczeniu spłaty zobowiązania w dyskretnych i skończonych modelach ewolucji terminowej struktury stóp procentowych.

ANALIZA PROJEKTÓW

Rozważania przedstawione w pracy *Dwukryterialna analiza projektu modelowanego z uwzględnieniem buforów czasowych i kosztowych* (P. Błaszczyk, M.B. Kania, T. Błaszczyk) wynikają z próby budowy modelu pozwalającego na poszukiwanie optymalnych rozwiązań problemu czasowo-kosztowego w projekcie planowanym na podstawie założenia niedokładności oszacowań czasu trwania i kosztów realizacji czynności. Przyjęto założenie, że przedkładane przez zainteresowanych wykonawców wielkości są zazwyczaj obciążone pewnym przeszacowaniem wynikającym z odpowiedzialności za późniejszą realizację czynności, zgodnie z przyjętym harmonogramem i budżetem. Zaproponowany w pracy model zakłada możliwość zachęcenia wykonawców do partycypacji w ryzyku niezrealizowania mniej ostrożnych szacunków w zamian za udział w korzyściach wynikających z ewentualnej szybszej i tańszej realizacji. Koncepcja teoretyczna zilustrowana jest przykładem liczbowym.

Praca *Zrównoważone podejście w zarządzaniu ryzykiem projektów europejskich* (M. Fedyczak, D. Kuchta) zawiera koncepcję zrównoważonej metody zarządzania ryzykiem publicznych projektów europejskich uwzględniającą identyfikację, analizę oraz reakcję na ryzyko w wielu płaszczyznach. Podejście stanowi próbę przeniesienia idei zrównoważenia i wielopłaszczyznowej analizy stanowiącej podstawę Zrównoważonej Karty Wyników na zarządzanie ryzykiem publicznych projektów europejskich. Autorki proponują perspektywy oceny i kontroli ryzyka w projektach europejskich, a także wskazują – na podstawie analizy studium przypadku – źródła ryzyka oraz błędy i problemy w trakcie realizacji, które mogą przyczynić się do porażki projektu (rozumianej głównie jako opóźnienie lub przekroczenie budżetu) w odniesieniu do poszczególnych perspektyw i faz projektu. Opracowanie zawiera także zarys dalszego rozwoju koncepcji zrównoważonego podejścia.

W pracy ***Nowa koncepcja odporności planów projektów*** (D. Kuchta) przedstawiono przegląd różnych podejść do odporności planów projektowych. Następnie zaprezentowano metodę zapewniania odporności systemom w innych dziedzinach, takich jak system przyjmowania zgłoszeń o awariach czy wypadkach. Metoda ta zakłada, że system będzie odporny wtedy, kiedy w żadnym momencie wymagana intensywność obsługi nie będzie wyższa niż ta, którą dany system jest w stanie zapewnić. Metoda ta zakłada również, że przy hierarchizowaniu problemów, z jakimi sobie trzeba radzić przy obsłudze systemu, brana jest pod uwagę zarówno waga problemu, jak i czas do awarii, czyli ilość czasu, jaka minie, zanim danego problemu nie da się już rozwiązać. Zaproponowano sposób przeniesienia tej koncepcji na planowanie projektów. Intensywność obsługi jest pochodną liczby i wagi problemów, jakie trzeba rozwiązać w związku z realizowanymi w danym momencie zadaniami, a czas do awarii to czas, jaki upłynie do momentu, kiedy dany problem nieodwołalnie spowoduje porażkę projektu (np. nieprzyjęcie go przez klienta). Zaproponowano model matematyczny oraz algorytm pozwalające wyznaczyć odporne plany projektu zgodne z tą koncepcją. Przedstawiono również przykład.

W pracy ***Zarządzanie ryzykiem poprzez uczenie się w projektach europejskich*** (D. Kuchta, E. Ptaszyńska) zaprezentowano koncepcję zarządzania ryzykiem przez uczenie się w projektach współfinansowanych przez Unię Europejską, a realizowanych przez jednostki samorządu terytorialnego. W pierwszej części pracy opisano projekty europejskie i ryzyko z nimi związane. Następnie przedstawiono koncepcję zarządzania ryzykiem przez uczenie się (znaną z literatury w odniesieniu do projektów z innej branży) oraz propozycję modyfikacji tej koncepcji dostosowaną do projektów europejskich. Zaprezentowano również przykład zastosowania opisanej koncepcji.

Jednym z dziewięciu obszarów zdefiniowanych w PMBOK Guide'2004 jest obszar Zarządzania Ryzykiem. Wspomniany podręcznik omawia między innymi narzędzia i techniki ilościowej analizy ryzyka. W pracy ***Wykorzystanie opcji realnych w szacowaniu ryzyka przekroczenia przez koszty projektu ustalonej wartości*** (K.S. Targiel) zaproponowano zastosowanie opcji realnych w zarządzaniu ryzykiem w projekcie jako narzędzia w ilościowej analizie ryzyka projektu. Bazując na porównaniu kosztów planowanych z poniesionymi, w rozważanym teoretycznym projekcie informatycznym oszacowano parametry procesu stochastycznego opisującego te różnice. Opierając się na konstrukcji opcji sprzedaży oszacowano parametry rozkładu prawdopodobieństwa rzeczy-

wistego kosztu w momencie zakończenia projektu. Znając taki rozkład, możliwe jest określenie prawdopodobieństwa przekroczenia przez projekt dopuszczalnego pułapu kosztów.

Zwinne zarządzanie projektami, mimo że narodziło się niedawno, zyskało sobie znaczącą popularność, szczególnie w branży oprogramowania. Za interesowanie tym podejściem wciąż wzrasta i coraz więcej firm, w tym dużych korporacji, wprowadza metodyki zwinne lub ich elementy do zarządzania projektami. Jedną z zalet metodyk lekkich jest minimalizacja negatywnych wpływów zmian otoczenia, w jakim osadzony jest projekt; obejmuje to także zmieniające się potrzeby i oczekiwania klienta. Zatem w metodyki te wbudowane są pewne mechanizmy pozwalające na eliminację lub obniżenie oddziaływania różnego rodzaju zagrożeń. Praca ***Zarządzanie ryzykiem w zwinnych metodykach zarządzania projektami*** (W. Walczak) ukazuje, w jaki sposób i w jakim zakresie metodyki lekkie realizują zarządzanie ryzykiem. Jednocześnie zidentyfikowane zostały te obszary zarządzania ryzykiem, które są zaniedbywane w istniejących metodykach zwinnych lub do których wskazane jest wprowadzenie usprawnień. Na tej podstawie zbudowano prezentowaną w opracowaniu koncepcję metody zarządzania ryzykiem, która może być stosowana wspólnie z najpopularniejszymi zwinnymi metodykami zarządzania projektami, stanowiąc ich uzupełnienie.

TEORIA GIER I NEGOCJACJI

Z literatury (przedmiotu) wiadomo, że przy wyznaczaniu oceny grupowej bądź wyniku głosowania rezultat w istotny sposób zależy od zastosowanej metody. Stwierdzenie to uzasadnia konieczność dokonania analizy zagadnienia na ile otrzymany wynik odzwierciedla opinie ekspertów (wolę wyborców) czy raczej cechy zastosowanej metody agregacji ocen ekspertów. Jeżeli nie ma narzuconych ograniczeń odnośnie wyboru metody oceny grupowej, w praktyce może okazać się, że rozwiązania uzyskane przy użyciu różnych metod są zgodne lub zbliżone albo całkowicie rozbieżne. W pierwszej sytuacji można przyjąć, że ocena grupowa wiernie odzwierciedla opinie ekspertów (wolę wyborców). W szczególnym przypadku różne metody mogą wskazywać ten sam obiekt jako zwycięzcę. Baharad i Nitzan podali warunki, jakie powinny być spełnione, aby zapewnić taki rezultat. Saari przedstawił metodologię umożliwiającą wyznaczanie wszystkich możliwych postaci ocen grupowych – odpowiadających różnym metodom pozycyjnym – dla danego zestawu ocen ekspertów oraz przeanalizo-

wał przykład uporządkowań trzech obiektów. W pracy *Wpływ wyboru metody wyznaczania oceny grupowej na wynik ekspertyzy* (H. Bury, D. Wagner) rozszerzono rozważania na przykład czterech obiektów oraz zwrócono uwagę na trudności analizy uzyskanych wyników przy większej liczbie obiektów.

W pracy *Znaczenie Senatu w polskim systemie legislacyjnym – przestrzenny model teoriogrowy* (R. Golański) rozpatrywany jest model grupowego podejmowania decyzji, w którym homogeniczne partie muszą dokonać wyboru określonej polityki w systemie stylizowanym na polski system parlamentarny. W izbie niższej każda z partii może zgłosić propozycję projektu odpowiadającą rzeczywistej pozycji zajmowanej przez daną partię lub (strategicznie) powstrzymać się od jej zgłoszenia. Jedna z partii ma charakter dominujący – może utworzyć koalicję większościową z każdą z pozostałych partii oraz kontroluje stanowisko marszałka izby niższej, który decyduje o kolejności poddawania projektów pod głosowanie. Partie wyłaniają zwycięski projekt, głosując strategicznie za lub przeciw rozpatrywanej propozycji. Projekt trafia następnie do izby wyższej, której preferencje odpowiadają preferencjom partii dominującej, a która może przedstawić projekt konkurencyjny – w takim przypadku projekt trafia z powrotem do izby niższej, gdzie jest przez nią ponownie rozpatrywany. Do rozwiązania gry zastosowano pojęcie równowagi doskonałej i rozwiązywalności przez dominację. Celem pracy jest zbadanie wpływu istnienia Senatu na przyjmowane projekty. W pracy przedstawiono przykład ilustrujący, że gdy nie istnieje zwycięzca w sensie Condorceta, to dla partii dominującej niekorzystne może być kontrolowanie stanowiska marszałka (inne partie mogą się wówczas strategicznie powstrzymać od zgłaszania określonych projektów). Pokazano również, że w grze, w której Senat pełni rolę legislacyjną wynik może być inny niż w grze, w której projekt przyjmowany jest przez parlament jednoizbowy.

W pracy *Ograniczanie ryzyka kosztów obsługi długu publicznego przy wykorzystaniu rozwiązania minimaxowego gry strategicznej* (L. Klukowski) przedstawiono propozycję metody konstrukcji portfela skarbowych instrumentów dłużnych, zapewniającego minimalizację ryzyka przyrostu kosztów obsługi długu, w przypadku wariantowej prognozy stóp procentowych. Proponowana koncepcja polega na zastosowaniu rozwiązania minimaxowego dwuosobowej gry strategicznej, ze skończoną liczbą strategii, o sumie zero, w której graczami są emitent i rynek finansowy (inwestorzy), a macierz gry wyraża przyrosty kosztów obsługi, wynikające z błędnego wyboru wariantu.

Wcześniejsza praca autorki *Ryzyko porządkowe a rachunek zdarzeń* zawiera wyniki badania, czy subiektywne ryzyko porządkowe rządzi się prawami rachunku prawdopodobieństwa. Wynik okazał się negatywny, zwłaszcza gdy chodzi o prawdopodobieństwo sumy zdarzeń. Badane było ryzyko zdarzeń negatywnych na przykładzie zachorowań na różne choroby. W pracy *Reguły formalne ryzyka porządkowego w przypadku zdarzeń negatywnych i pozytywnych* (H. Sosnowska) to badanie powtórzone zostało w przypadku wybranych zdarzeń pozytywnych (takich jak dobre ukończenie studiów czy wysokie zarobki). Ponownie wyniki nie są zgodne z regułami rachunku prawdopodobieństwa, przy czym ta niezgodność jest większa niż w przypadku zdarzeń negatywnych. Porównane zostały wyniki obu badań.

Praca *Quasi-klasyczne metody wspomagania racjonalnego podziału wygranej w kooperacyjnych grach z koalicjami* (M. Wolny) prezentuje metody (wspomagania) podziału wygranej w grach koalicyjnych, w których istnieje niepusty rdzeń, czyli możliwy jest przynajmniej jeden podział wygranej spełniający jednocześnie postulaty racjonalności: zbiorowej, indywidualnej i koalicyjnej. Obok klasycznych metod podziału wygranej (wartość Shapleya, wartość Banzhafa, Nucleolus, punkt Gately'ego) przedstawiono koncepcje wspomagania podziału wygranej wykorzystujące metody wielokryterialne: programowanie leksykograficzne, metodę sumy ważonej, metody punktu referencyjnego oraz programowania celowego. Metody wielokryterialne uznane zostały za quasi-klasyczne dla problemu określenia podziału wygranej. Uzasadnieniem stosowania metod wielokryterialnych może być sytuacja, gdy zaakceptowany, uznany za sprawiedliwy, sposób podziału wygranej między graczy implikuje podział, który nie jest racjonalny – podział ten jest poza rdzeniem gry. Zaprezentowane koncepcje przedstawiono na przykładzie gry przedsiębiorstw o rynek.

Praca *Równowaga w grze fiskalno-monetarnej a priorytety banku centralnego i rządu* (I. Woroniecka-Leciejewicz) przedstawia analizę dwuosobowej gry między bankiem centralnym a rządem, zwanej grą fiskalno-monetarną, ze skończoną liczbą strategii w zakresie polityki pieniężnej i budżetowej, różniących się stopniem ich restrykcyjności/ekspansywności. Przeprowadzono analizę stanów równowagi i Pareto-optymalności rozwiązań w takiej grze w przypadku alternatywnych założeń dotyczących uwarunkowań koniunktury gospodarczej i skuteczności polityki monetarnej i fiskalnej, a także z uwzględnieniem różnych priorytetów banku centralnego i rządu w kształtowaniu poli-

tyki makroekonomicznej. Odmienne założenia znajdują odzwierciedlenie w inaczej zdefiniowanych wypłatach w grze, a w konsekwencji w stanach równowagi gry i ich Pareto-optymalności. W zależności od przyjętych założeń dotyczących uwarunkowań koniunktury gospodarczej, przede wszystkim wpływu deficytu budżetowego na wzrost PKB, niezależnie działające władze monetarne i fiskalne dążą, zgodnie z równowagą Nasha, do wyboru restrykcyjnej polityki pieniężnej i ekspansywnej budżetowej (wariant A) bądź do wyboru obu restrykcyjnych rodzajów polityk (wariant B). Ponadto, podjęto próbę analizy powyższej gry monetarno-fiskalnej rezygnując z założenia o liniowych zależnościach między miernikami stanu gospodarki (wzrostu gospodarczego i inflacji) a instrumentami policy mix: deficytem budżetowym w relacji do PKB i realną stopą procentową. Wykorzystując twierdzenie i wzór Taylora, wyprowadzono formuły na wartości występujące w tablicy wypłat w grze dla nieliniowych zależności, a następnie przeprowadzono, analogicznie jak dla zależności liniowych, analizę stanów równowagi gry i ich Pareto-optymalności.

Tadeusz Trzaskalik

OPTYMALNE DECYZJE

ZASTOSOWANIE MODELI PROGRAMOWANIA DYNAMICZNEGO W WARUNKACH ROZMYTYCH W ZARZĄDZANIU MAJĄTKIEM

Wprowadzenie

We wcześniejszych pracach autora [1; 2] zaproponowano wykorzystanie modeli programowania dynamicznego w warunkach rozmytych we wspomaganiu planowania inwestycyjnego – do wyboru oraz szeregowania projektów inwestycyjnych. W proponowanym opracowaniu rozważane jest szersze wykorzystanie tego typu modelowania w zarządzaniu majątkiem firmy lub jednostki samorządowej – rozpatrywane są nowe grupy celów, a oprócz realizacji projektów inwestycyjnych jako poszczególne zadania występują także operacje zbycia elementów majątku. Proponowane podejście do konstruowania planu inwestycyjnego (planu zarządzania majątkiem) zainspirowane zostało pracami [3; 4] oraz [5], wykorzystywane metody bazują na podanych w [4], szczególne zaś wersje algorytmu ewolucyjnego zaproponowano w [6].

1. Model decyzyjny

Zadanie dotyczy ustalenia kolejności podejmowania zadań inwestycyjnych z listy zadań do wykonania (alternatywnie: lista może być dłuższa, a ustalenie kolejności ma wówczas jednocześnie pozwolić wybrać projekty do realizacji – w takiej wersji zadania nie realizowane projekty znajdują się na końcu listy, poza horyzontem planowania).

Wyjściowy model (M1L) formułujemy następująco

$$R(\mathbf{x}) = \tilde{g}_1(x_1) + \tilde{g}_2(x_2) + \dots + \tilde{g}_N(x_N) \rightarrow \max \quad (1)$$

$$\sum_{i=1}^N a_i(x_i) = q \quad (2)$$

$$x_i \in \{1, 2, \dots, N\}, \quad x_i \neq x_j \text{ dla } i \neq j \quad (3)$$

gdzie:

$\tilde{g}_i(x_i)$ – korzyść z projektu x_i , tj. podjętego na etapie i -tym zadania inwestycyjnego, może to być również zagregowana ocena jednoczesnej realizacji kilku celów (w specyfikacji modelu można również ująć każdy z rozpatrywanych celów z osobna), dana liczbą rozmytą (lub wartością lingwistyczną),

$a_i(x_i)$ – koszt zadania x_i (może być zmienny i zależny od etapu, stąd indeks przy a),

q – łączny budżet na inwestycje w całym okresie,

N – liczba etapów – zadań inwestycyjnych.

Możemy także wprowadzić alokację dwóch zasobów (drugi to np. własny sprzęt wykorzystywany przy podejmowanych inwestycjach)

$$\sum_{i=1}^N b_i(x_i) \leq s \quad (4)$$

Podstawowe równanie ujmujące proces alokacji, a tym samym definiujące funkcję przejścia jest postaci

$$f_N(q) = \max_{a_N(x_N)} (\tilde{g}_N(x_N) + f_{N-1}(q - a_N(x_N))) \quad (5)$$

Jeśli koszty realizacji nie zależą od etapu, wówczas warunek 2 jest spełniony z definicji (bilansuje wykorzystanie środków) i można go pominąć, a równanie (5) zastąpić równaniem

$$f_N(\mathfrak{X}) = \max_{x_N} (\tilde{g}_N(x_N) + f_{N-1}(\mathfrak{X} \setminus \{x_N\})) \quad (6)$$

gdzie $\mathfrak{X}$ – zbiór zadań do wykonania, $x_N \in \mathfrak{X}$.

Tak skonstruowany model należy uzupełnić o warunki dotyczące następstw w czasie poszczególnych zadań podejmowanych do realizacji, zapiszemy je w postaci

$$k < l, \text{ dla } (x_k, x_l) \in P \quad (7)$$

tzn. dla takich par (x_k, x_l) , dla których określona jest relacja P poprzedzania zadania x_l przez zadanie x_k , w rozwiązaniu ma zachodzić $k < l$.

Zapisane zadanie należy przeformułować w celu rozwiązania go algorytmem programowania dynamicznego – określamy etapy, zbiór stanów dopuszczalnych na każdym z etapów, zbiór decyzji dopuszczalnych dla każdego ze stanów, czyli funkcję przejścia* (daną pierwotnie wzorem 5), a także funkcje oceny: etapowe oraz łączną.

W rozszerzonej wersji modelu dopuszczalne jest wprowadzenie wyboru między alternatywnymi względem siebie wariantami każdego z podejmowanych zadań. Łączny budżet jest wtedy większy niż możliwości finansowe, a ostatnie z listy zadanie (zadania) nie będą realizowane, stanowiąc jedynie formalny wpis w planie inwestycyjnym. W zapisie modelu w miejscu wzoru (3) pojawiają się wówczas wzory (8)-(9)

$$(x_i, x_j) \notin I \text{ dla } i, j \in \{1, 2, \dots, N\} \quad (8)$$

$$x_i \in \{1, 2, \dots, M\}, M \geq N, x_i \neq x_j \text{ dla } i \neq j \quad (9)$$

gdzie:

I – relacja, w której są pary różnych wariantów tego samego przedsięwzięcia lub po dwie wzajemnie wykluczające się inwestycje.

Zauważmy, że współczynniki a_i , dotyczące kosztów podejmowanych do realizacji zadań, nie muszą być dodatnie, tzn. dopuszczalne jest, aby zadania podejmowane do realizacji powiększały, a nie pomniejszały pozostałą do wykorzystania część budżetu inwestycyjnego. Stąd propozycja rozszerzenia wykorzystania modelu na przypadek, w którym firma – w ramach zarządzania swoim majątkiem – część środków na inwestycje może pozyskać ze sprzedaży niewykorzystywanych składników majątku. Takie – nieinwestycyjne – zadania służyć mogą realizacji innych celów, niż dotychczas rozpatrywane (efekty ekonomiczne, efekty ekologiczne i zadowolenie pracowników oraz okolicznych mieszkańców), w szczególności możliwe jest wyróżnienie (i dołączenie ich do grupy jako celów kolejnych bądź inkorporacja w ramach niektórych z celów już istniejących) następujących: poprawa struktury organizacji, wzrost efektywności zarządzania, polepszenie płynności finansowej i warunków realizacji przedsięwzięć o charakterze inwestycyjnym (między innymi zwiększenie zdolności kredytowej). Dwa pierwsze z tej grupy celów można dołączyć do grupy poprawy warunków pracy, pozostałe natomiast w zasadzie służą szeroko rozumianej efektywności ekonomicznej – możliwa jest więc ocena nowowprowadzanych elementów planu zarządzania majątkiem w ramach poprzednio ustalonej wiązki celów.

* Dynamizacji zadania można dokonać w sposób intensywny bądź ekstensywny, tzn. wykorzystując warunki ograniczające bezpośrednio w ramach procedury określania stanów i decyzji dopuszczalnych na każdym etapie, bądź też określając odpowiednie zbiory stanów i decyzji z wybiórczym uwzględnieniem warunków ograniczających, niektóre z nich w odpowiedni sposób włączając do sformułowania funkcji celu.

2. Przykłady liczbowe

Niżej zaprezentowano trzy przykłady liczbowe dotyczące zastosowania opisywanych modeli. Przykład pierwszy dotyczy wykorzystania modelu M1L, przykład drugi to jego zmodyfikowana wersja, uwzględniająca zarówno możliwość sprzedaży elementów majątku, jak i występowanie wzajemnie wykluczających się zadań na liście (z których tylko jedno jest realizowane, drugie zaś zostaje umieszczone na końcu listy i nie jest przyjmowane do realizacji). Przykład trzeci to, zaczerpnięty z pracy [1], rzeczywistych rozmiarów przykład wykorzystania jednego z modeli do wygenerowania planu inwestycyjnego. W tym ostatnim przypadku do uzyskania rozwiązania (suboptymalnego) posłużono się algorytmem ewolucyjnym z krzyżowaniem OX, opisywanym między innymi w pracach [6] oraz [1].

Przykład 1

W przykładzie pierwszym należy uszeregować 9 zadań inwestycyjnych. Nazwy zadań oraz warunki, jakie ma spełniać ich lista, podano w tabeli 1. Wartości realizacji poszczególnych grup celów przez każde zadanie w zależności od pozycji na liście zawierają tabele 2-4. Wykorzystany model to M1L. Rozwiązanie znaleziono metodą programowania dynamicznego.

W wyniku obliczeń uzyskano dwie ścieżki optymalne – (9, 1, 5, 3, 2, 8, 4, 7, 6) oraz (9, 1, 5, 3, 2, 8, 4, 6, 7).

Tabela 1

Lista zadań inwestycyjnych i warunków narzuconych na permutacje

Nr	Zadanie	Warunki	Warunki
1	G-DO-4 Modernizacja stacji Dolna Odra 220/110/15 kV	Przed 3 i 4	–
2	G-DO-6 Budowa instalacji katalitycznego odazotowania spalin	Po 3	Przed 4 lub 8
3	G-DO-7-A Budowa nowego bloku gazowo-parowego	Przed 4	–
4	G-DO-7-B Budowa nowego bloku gazowo-parowego	Po 3	–
5	G-DO-8 Budowa portu barkowego	Przed 3 i 4	–
6	W-DO-1 Modernizacja IOS	–	–
7	W-DO-3 Modernizacja elektrofiltrów	–	–
8	W-SZ-1 Budowa kotła na biomasę	Po 9	Przed 1 lub 4
9	W-SZ-4 Przystosowanie turbiny do nowych parametrów pracy	Przed 8	Przed 3 lub 4

Tabela 2

Stopień realizacji celu *efekt ekonomiczny* w zależności od pozycji na liście,
dany wartościami lingwistycznymi

Nr	Pozycja zadania na liście								
	1	2	3	4	5	6	7	8	9
1	duży	duży	duży	średni++	średni++	średni	średni	średni	średni
2	duży	duży	duży	duży	duży	średni++	średni++	średni++	średni++
3	wielki	wielki	wielki	wielki-	b.duży++	duży	duży	duży	duży
4	b. duży	b. duży	wielki	wielki	wielki	b.duży++	b. duży	duży	duży
5	wielki	wielki	b.duży++	b.duży++	b. duży	b. duży	duży++	duży	duży
6	średni++	średni++	średni++	duży-	duży	duży+	duży+	duży++	duży++
7	średni++	średni++	średni++	duży-	duży	duży+	duży+	duży++	duży++
8	wielki	wielki	wielki	wielki	wielki-	b.duży++	b. duży+	b. duży	duży++
9	wielki	wielki	wielki	wielki-	b.duży++	b. duży+	b. duży	duży++	duży++

Źródło: Na podstawie [1, s. 112].

Tabela 3

Stopień realizacji celu *bhp* w zależności od pozycji na liście,
dany wartościami lingwistycznymi

Nr	Pozycja zadania na liście								
	1	2	3	4	5	6	7	8	9
1	średni	średni	średni	średni-	średni-	średni-	mały++	mały+	mały+
2	duży	duży	duży	duży	duży	średni++	średni++	średni++	średni++
3	duży++	duży++	duży++	duży	duży	średni++	średni++	średni++	średni++
4	duży++	duży++	duży++	duży	duży	średni++	średni++	średni++	średni++
5	średni++	średni++	średni++	średni	średni	średni	średni	średni	średni
6	mały	mały	mały	mały	mały	mały	mały	mały	mały
7	mały	mały	mały	mały	mały	mały	mały	mały	mały
8	duży++	duży++	duży++	duży++	duży	duży	duży	duży	duży
9	duży++	duży++	duży++	duży++	duży	duży	duży	duży	duży

Tabela 4

Stopień realizacji celu *ochrona środowiska* w zależności od pozycji na liście,
dany wartościami lingwistycznymi

Nr	Pozycja zadania na liście								
	1	2	3	4	5	6	7	8	9
1	średni	średni	średni	średni	średni-	średni-	średni-	średni-	średni-
2	b.duży++	b.duży++	b.duży++	b.duży+	b.duży	b.duży-	duży++	duży	duży
3	b.duży++	b.duży+	b.duży	b.duży	b.duży-	b.duży-	b.duży-	duży++	duży++
4	b. duży	b. duży	b.duży++	b.duży+	b.duży	b.duży-	b.duży-	duży++	duży++
5	średni(-)	średni(-)	mały(-)	mały(-)	mały(-)	mały(-)	mały(-)	mały(-)	mały(-)
6	b.duży++	b.duży++	b.duży+	b.duży+	b.duży	b.duży	b.duży-	b.duży-	b.duży-
7	duży++	duży++	duży++	duży++	duży+	duży+	duży+	duży	duży
8	wielki	wielki	wielki	wielki-	b.duży++	b.duży+	b.duży	b.duży-	duży++
9	wielki	wielki	wielki-	b.duży++	b.duży++	b.duży	b.duży-	duży++	duży++

Przykład 2

Przykład drugi jest rozszerzeniem przykładu poprzedniego o dwie operacje zbycia nieruchomości należących do przedsiębiorstwa – w miejsce zadań 3 i 4 w poprzednim przykładzie pojawiają się teraz trzy zadania N-DO-1, N-DO-2, N-DO-2' (ostatnie z nich jest wariantem zadania N-DO-2, dającym przychody w ratach). Ponieważ lista zadań składa się teraz z 10 elementów, dlatego dla elementów listy z przykładu 1 dokonano przepisania wartości ocen realizacji wszystkich celów z pozycji (kolumny) dziewiątej do kolumny dziesiątej. Na nowe zadania narzucono warunki jak w tabeli 5 oraz przypisano wartości ocen realizacji dwóch grup celów (*efekt ekonomiczny* i *bhp*) podane w tabelach 6 i 7.

Uzyskano rozwiązania optymalne dane ścieżkami (9, 10, 8, 1, 5, 11, 2, 7, 6, 12) oraz (9, 10, 8, 1, 5, 11, 2, 6, 7, 12).

Tabela 5

Lista zadań inwestycyjnych i warunków narzuconych na permutacje

Nr	Zadanie	Warunki	Warunki
10	N-DO-1 Zbycie nieruchomości nr 1	Przed 2 i 8	–
11	N-DO-2 Zbycie nieruchomości nr 2	Przed 2	Po 10
12	N-DO-2' Zbycie nieruchomości nr 2 (drugi wariant)	Przed 2	Po 10

Tabela 6

Stopień realizacji celu *efekt ekonomiczny* w zależności od pozycji na liście, dany wartościami lingwistycznymi

Nr	Pozycja zadania na liście									
	1	2	3	4	5	6	7	8	9	10
10	duży+	duży+	duży+	duży	duży	duży	duży-	duży-	duży-	duży-
11	duży	duży	duży	duży	duży-	duży-	średni+	średni+	średni	średni
12	duży-	duży-	średni+	średni+	średni	średni	średni+	średni	średni	średni

Tabela 7

Stopień realizacji celu *bhp* w zależności od pozycji na liście, dany wartościami lingwistycznymi

Nr	Pozycja zadania na liście									
	1	2	3	4	5	6	7	8	9	10
10	średni	średni	średni	średni	średni-	średni-	średni-	średni-	średni-	średni-
11	średni-	średni-	średni-	średni-	średni	średni	średni	średni+	średni+	średni+
12	średni-	średni-	średni-	średni-	średni	średni	średni	średni+	średni+	średni+

Przykład 3

Przykład 3 zaczerpnięty został z wcześniejszej pracy autora [1, s. 114], w której dokonano między innymi specyfikacji rodziny modeli z lingwistycznymi wartościami ocen spełnienia celów oraz warunków zbliżonych do M1L (model został uzupełniony między innymi o warunki związane z latami, wskutek czego nie jest możliwe rozwiązywanie zadań metodą programowania dynamicznego wstecz – szerzej na ten temat w pracy [1]).

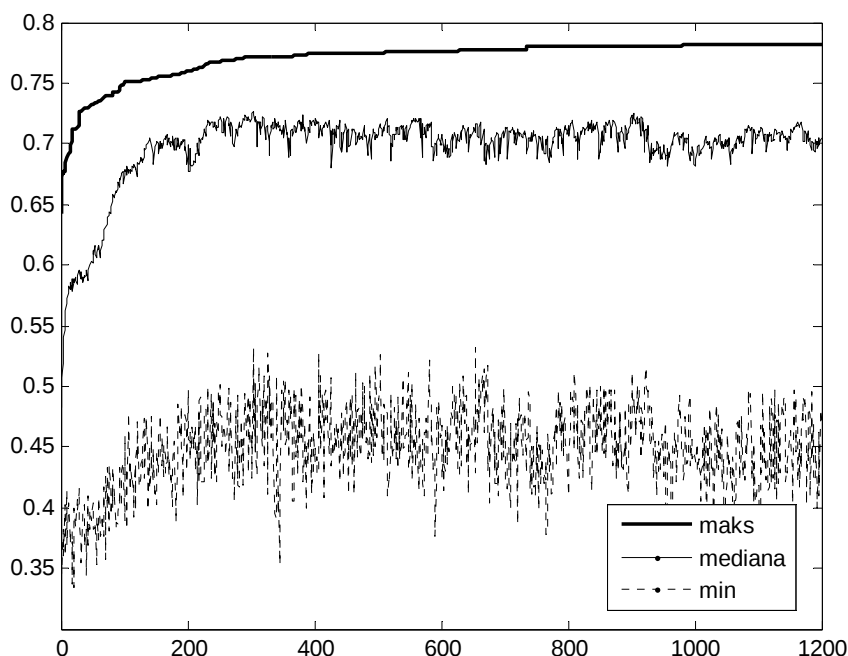
Rozwiązanie zadania uzyskano z wykorzystaniem algorytmu ewolucyjnego dla zagadnień permutacji, z krzyżowaniem OX i strategią elitarystyczną [6; 1]. Rysunek 1 przedstawia proces poprawy rozwiązania – na osi poziomej odłożone są numery kolejne pokoleń (iteracji), na osi pionowej – wartości łącznej funkcji oceny rozwiązania na danym etapie najlepszego, przeciętnego (mediana funkcji ocen) oraz najłabszego. Wartości najlepsze ocen nie maleją, gdyż zastosowano strategię elitarystyczną (najlepsze osobniki w populacji przechodzą do populacji potomnej).

Tabela 8

Rozwiązanie suboptymalne zadania z przykładu 3 – 1200 pokoleń, 500 osobników, 8 osobników elity, OX, lingwistyczne wartości ocen

Rozwiązanie	134	107	94	26	80	4	133	126	10	92
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	79	99	29	117	11	72	68	70	132	120
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	40	41	102	103	2	71	77	121	130	50
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	136	15	17	137	12	45	13	106	135	16
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	131	23	37	34	109	65	48	63	53	66
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	49	56	113	100	127	95	114	105	129	98
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	90	101	73	82	81	14	51	97	123	83
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	111	104	19	24	118	67	64	46	52	110
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	28	3	59	84	122	47	119	115	62	54
	2008	2008	2008	2008	2008	2008	2008	2008	2008	2008
	93	78	20	124	8	42	6	86	89	116
	2008	2008	2008	2008	2008	2009	2009	2010	2010	2010
	88	36	87	5	27	38	60	108	74	35
	2011	2012	2012	2012	2013	2013	2013	2013	2013	2013
	44	128	112	58	55	39	61	18	43	33
	2013	2013	2013	2013	2013	2013	2013	2013	2013	2013
	25	32	125	57	75	96	21	31	69	91
	2013	2013	2013	2013	2013	2013	2013	2013	2013	2013
	9	76	22	30	7	85	1			Ocena
	2013	2013	2013	2013	2013	2014	2014			~0,77

Źródło: [1, s. 142].



Rys. 1. Poprawa rozwiązania w przykładzie 3, 1200 pokoleń, 500 osobników, 8 osobników elity, OX, lingwistyczne wartości ocen

Źródło: [1, s. 115].

Podsumowanie i wnioski

Przeprowadzone badanie pozwala wysnuć ogólny wniosek, że możliwe jest skuteczne modelowanie planu inwestycyjnego (planu zarządzania majątkiem) w warunkach niepełnej informacji, czego wyrazem są uzyskane wyniki. Interpretacja tak uzyskanych wyników modelowania nie sprawia żadnej trudności, trudne może być jedynie przeniesienie wiedzy w odpowiedniej postaci na sformułowany model i pokonanie bariery, głównie psychologicznej natury, uniemożliwiającej korzystanie z takich, względnie nowych, metod. Wprowadzenie modyfikacji w postaci zadań sprzedaży elementów majątku daje się łatwo zaimplementować i nie ma negatywnego wpływu ani na model, ani na metodę rozwiązywania, daje także sensowne rezultaty. Jednocześnie umożliwia to wprowadzenie do modelu nowych celów, które mogą być włączone w skład wcześniej wyspecyfikowanych grup celów (jak to zostało pokazane w przykładzie 2) bądź też stanowić cele odrębne. Dodatkowo dla zadań rzeczywistych

rozmiarów, dla których programowanie dynamiczne jest zawodne, istnieją skutecznie działające algorytmy znajdowania rozwiązań suboptymalnych (przykład 3).

Literatura

1. Baran P. (2008). Rozmytość II typu w dynamicznym programowaniu inwestycji. W: Modelowanie preferencji a ryzyko '08. Red. T. Trzaskalik. AE, Katowice.
2. Baran P. (2010). Wybrane metody programowania dynamicznego we wspomaganiu decyzji inwestycyjnych przedsiębiorstwa sektora energetycznego. Uniwersytet Szczeciński, Szczecin.
3. Bellman R.E., Zadeh L.A. (1970). Decision Making in a Fuzzy Environment. *Management Science*, 4 (17).
4. Kacprzyk J. (2001). Wieloetapowe sterowanie rozmyte. Wydawnictwo Naukowo-Techniczne, Warszawa.
5. Trzaskalik T., Sitarz S. (2005). Modele programowania dynamicznego w strukturach porządkowych. W: *Badania Operacyjne i Systemowe 2004. Podejmowanie decyzji*. Instytut Badań Systemowych PAN, Akademicka Oficyna Wydawnicza EXIT, Warszawa.
6. Whitley D. (1994). A Genetic Algorithm Tutorial. *Statistics and Computing*, 4.

Krzysztof Dmytrów

Uniwersytet Szczeciński

ZASTOSOWANIE MATEMATYCZNEGO MODELU ZAPASÓW DLA PRODUKTÓW PSUJĄCYCH SIĘ DO GOSPODAROWANIA ZASOBAMI GOTÓWKI

Wstęp

Modelowanie decyzyjne gospodarowania zapasami materiałowymi może być wykorzystane nie tylko do kontrolowania zakupów i zapasów dokładnie takich pozycji. Okazuje się, że podobny sposób postępowania można zastosować w gospodarowaniu gotówką w banku. Podobnie jak w przypadku zapasów materiałowych, tutaj także mamy do czynienia z dostawą produktu (w przypadku banku będzie to np. zasilenie bankomatu), przechowywaniem czy niedoborami. Analogie występują także w popycie zgłaszanym przez klientów, zmniejszaniem się poziomu zapasów, złożeniem zamówienia, dostawą itd. Problem gospodarowania gotówką w bankomacie banku komercyjnego z wykorzystaniem modelu zapasów został przedstawiony w pracy [1]. Wydaje się jednak, że model przedstawiony w powyższej pracy nie opisuje ważnej właściwości gotówki, a mianowicie utraty jej wartości spowodowanej inflacją. W takim przypadku pomocne staje się wykorzystanie modeli gospodarowania zapasami dla produktów psujących się.

W literaturze zapasów psucie się produktów jest rozumiane różnie. Może to być utrata masy, wartości albo walorów użytkowych na skutek rozkładu, parowania, przeciekania itd. Bardzo obszerny przegląd literatury i metod stosowanych w gospodarowaniu zapasami psującymi się przedstawiono w pracy [3].

W niniejszej pracy zdecydowano się rozwinąć prosty, deterministyczny model gospodarowania psującymi się zapasami zaprezentowany w pracy [2] do stochastycznego modelu $\langle Q, r \rangle$ dla zaległych zamówień i stałej stopy psucia się. Klasyczny model $\langle Q, r \rangle$ gospodarowania zapasami dla zaległych zamówień został zaprezentowany np. w pracy [4]. Tak więc przedstawiony w niniejszej pracy model jest uogólnieniem przypadkiem modelu wspomnianego w poprzednim zdaniu. Osiągnięciem jest uwzględnienie faktu psucia się produktu i zastosowanie modelu do gospodarowania gotówką.

1. Założenia i model

W pracy przyjęto następujące oznaczenia:

- λ – średnie zapotrzebowanie na gotówkę w ciągu roku,
- Q – optymalna wielkość zasilenia bankomatu w gotówkę,
- r – poziom zamawiania (jeżeli poziom gotówki w bankomacie spadnie poniżej tego poziomu, składa się zamówienie uzupełniające),
- C – jednostkowy koszt gotówki,
- A – koszt jednego zasilenia bankomatu w gotówkę,
- h – jednostkowy koszt magazynowania gotówki,
- π – jednostkowy koszt niedoboru gotówki,
- τ – okres realizacji dostaw, zwany też okresem wyprzedzenia (jest to okres, po jakim od momentu złożenia zamówienia uzupełniającego gotówka znajdzie się w bankomacie),
- $\bar{\eta}(r)$ – oczekiwana wielkość niedoborów gotówki w ciągu jednego okresu odnowienia zapasów,
- φ – stopa psucia się produktu.

W modelu $\langle Q, r \rangle$ należy przyjąć następujące założenia [4, s. 162-163]:

- koszt jednostkowy C jest stały, niezależny od Q ,
- jednostkowy koszt niedoboru π nie zależy od długości czasu, w którym występują niedobory,
- występuje tylko jedno zamówienie przychodzące,
- koszty operacyjne systemu zarządzania gotówką są niezależne od Q i r ,
- poziom zamawiania r jest dodatni.

Dodatkowo założono, że stopa psucia się jest stała.

Funkcja oczekiwanych kosztów gospodarki materiałami w modelu $\langle Q, r \rangle$ dla zaległych zamówień jest następująca [4, s. 166]

$$K(Q, r) = \frac{\lambda}{Q} A + h \left(\frac{Q}{2} + r - \mu \right) + \frac{\pi \lambda}{Q} \bar{\eta}(r) \rightarrow \min \quad (1)$$

gdzie [6]

$$\bar{\eta}(r) = \int_r^{\infty} (x - r) f(x) dx \quad (2)$$

oraz gdzie $f(x)$ jest funkcją gęstości prawdopodobieństwa rozkładu zapotrzebowania na gotówkę w okresie realizacji dostaw.

Wyznaczenie optymalnych wielkości Q oraz r polega na wyznaczeniu pierwszych pochodnych równania (1) względem Q oraz r i przyrównaniu ich do zera

$$\frac{\partial K(Q, r)}{\partial Q} = \frac{\partial K(Q, r)}{\partial r} = 0$$

Tak więc optymalna wielkość Q będzie dana wzorem

$$Q = \sqrt{\frac{2\lambda[A + \pi\bar{\eta}(r)]}{h}} \quad (3)$$

Optymalną wielkość r natomiast otrzymamy za pomocą następującego wzoru

$$F(r) = 1 - \frac{hQ}{\pi\lambda} \quad (4)$$

gdzie $F(r) = \int_0^r f(x)dx$.

Powyższy model zakłada, że nie występuje psucie się produktu. Przyjmując teraz, że produkt ulega psuciu się i stopa psucia się jest stała w czasie, w ciągu jednego okresu (zakładamy, że jest to rok) zepsuje się stały odsetek ilości – φ . Dlatego w ciągu jednego cyklu odnowienia zapasów zepsuje się odsetek $\varphi \cdot 1/n$, gdzie n jest liczbą zamówień. Wiedząc, że liczba zamówień dana jest następującym wzorem

$$n = \frac{\lambda}{Q}$$

W takim przypadku jednorazowo kupi się [2, s. 423]

$$\left(1 + \varphi \frac{Q}{\lambda}\right)Q = Q^* \quad (5)$$

ażeby pokryć zapotrzebowanie na produkt.

Tak więc podstawiając Q^* z równania (5) za Q we wzorze (1) otrzymamy

$$K(Q, r) = \frac{\lambda}{\left(1 + \varphi \frac{Q}{\lambda}\right)Q} A + h \left[\frac{Q}{2} \left(1 + \varphi \frac{Q}{\lambda}\right) + r \left(1 + \varphi \frac{Q}{\lambda}\right) - \mu \right] + \frac{\pi\lambda}{\left(1 + \varphi \frac{Q}{\lambda}\right)Q} \bar{\eta}(r) \rightarrow \min \quad (6)$$

Różniczkując równanie (6) po Q otrzymamy

$$\frac{\partial K}{\partial Q} = -\frac{\lambda + 2\varphi Q}{\left(Q + \varphi \frac{Q^2}{\lambda}\right)^2} [A + \pi\bar{\eta}(r)] + \frac{h}{2} + h\varphi \frac{Q}{\lambda} + h\varphi \frac{r}{\lambda} = 0 \quad (7)$$

Ponieważ w równaniu (7) Q występuje w różnych potęgach, dlatego nie da się otrzymać zamkniętej formy rozwiązania i trzeba się uciec do metod numerycznych*.

Z kolei wyznaczenie pierwszej pochodnej względem r i przyrównanie jej do zera da taki sam wynik, jak równanie (4).

2. Przykład numeryczny

W pracy wykorzystano dane umowne. Jednostką zapotrzebowania było 100 zł, średnie zapotrzebowanie na rok określono na 1 000 000 zł, czyli w przyjętych jednostkach $\lambda = 10\,000$. Odchylenie standardowe zapotrzebowania (S) przyjęto na 500. Okresem realizacji dostaw (τ) był 1 tydzień, czyli około 1/53 roku. Średnie zapotrzebowanie w okresie realizacji dostaw oszacowano za pomocą wzoru: $\mu = \lambda\tau = 188,679$, zaś odchylenie standardowe za pomocą wzoru: $\sigma = S\sqrt{\tau} = 88,680$. Założono, że rozkład zapotrzebowania w okresie realizacji dostaw jest rozkładem normalnym o parametrach: $N(188,679; 88,680)$. Koszt zasilenia bankomatu (A) przyjęto na poziomie 150 zł, jednostkowy koszt magazynowania gotówki (h) wyniósł 6 zł w ciągu roku (przyjęto go na poziomie utraconych korzyści, jakie można osiągnąć trzymając gotówkę na rachunku oprocentowanym 6% w skali roku). Jednostkowy koszt niedoboru (π) przyjęto na poziomie 20 zł jako koszt zaciągnięcia kredytu w razie braku gotówki. Stopę psucia się przyjęto na poziomie od 0 do 0,5 co 0,1. Ponieważ za stopę psucia się przyjęto poziom inflacji, dlatego te liczby nie oddają sytuacji rzeczywistej. Tak duże wartości przyjęto po to, ażeby były widoczne różnice w porównaniu z przypadkiem, gdy psucie się (inflacja) nie występuje. Wyniki obliczeń przedstawia tabela 1.

Tabela 1

Wyniki obliczeń optymalnej wielkości zamówienia S^* , poziomu zamawiania (r^*), oczekiwanej ilości niedoborów ($\bar{\eta}(r)$) oraz oczekiwanych łącznych rocznych kosztów

φ	0	0,1	0,2	0,3	0,4	0,5
Q^*	733,5088	730,9481	728,4788	726,0943	723,7885	721,5560
r^*	308,1237	310,4808	312,7892	315,7908	318,2492	320,6634
$\bar{\eta}(r)$	1,1411	1,1260	1,1116	1,0978	1,0845	1,0717
$K(Q^*, r^*)$	5 117,72 zł	5 129,83 zł	5 141,78 zł	5 157,99 zł	5 171,05 zł	5 183,94 zł

* W niniejszej pracy posłużono się funkcją `fsolve` zawartą w pakiecie Scilab 5.1.

Dla parametru $\varphi = 0$ wyniki są identyczne z klasycznym przypadkiem modelu $\langle Q, r \rangle$ dla zaległych zamówień, gdy nie zachodzi psucie się produktów. Otrzymane wyniki pokazują, że wraz ze wzrostem stopy psucia się (w naszym przypadku inflacji), maleje wielkość zamówienia oraz rośnie poziom zamawiania. Wzrost poziomu zamawiania powoduje, że maleje oczekiwana ilość niedoborów w jednym cyklu zasilenia bankomatu w gotówkę. Im wyższa stopa psucia się, tym większe oczekiwane łączne roczne koszty. Wynika to z dwóch powodów. Po pierwsze, mniejsze wielkości zamówienia powodują, że bank zmuszony jest częściej uzupełniać gotówkę, co podnosi łączne koszty zamawiania. Po drugie, wyższy poziom zamawiania powoduje konieczność utrzymywania wyższego zapasu bezpieczeństwa, co podnosi łączne koszty magazynowania gotówki. Z kolei oczekiwane koszty niedoborów są coraz mniejsze wraz ze wzrostem stopy psucia się, jednak udział kosztów niedoboru w łącznych kosztach gospodarowania gotówką jest relatywnie niewielki.

Wnioski

W niniejszej pracy zaprezentowano prosty model $\langle Q, r \rangle$ dla zaległych zamówień uwzględniający psucie się produktów. Przykładem zastosowania takiego modelu może być gospodarowanie zasilaniem i przechowywaniem gotówki w bankomacie banku komercyjnego. W przypadku gotówki psucie się jest spowodowane utratą jej wartości, czyli inflacją. Otrzymane wyniki wskazują, że wraz ze wzrostem stopy inflacji łączne koszty gospodarowania gotówką wzrastają. Wzrost kosztów spowodowany jest wzrostem kosztów zamawiania i magazynowania gotówki. Widać jednak, że nawet przy założeniu wysokiej, bo 50% inflacji, wzrost oczekiwanych łącznych kosztów gospodarowania zapasami w porównaniu z sytuacją, w której inflacja nie wystąpiłaby, wynosi niecałe 70 zł. Na tle łącznych kosztów wynoszących ponad 5000 zł, stanowi to niecałe 1,5%. W praktyce uwzględnianie stopy inflacji przy zastosowaniu modelu gospodarowania zapasami w gospodarowaniu gotówką miałoby sens dla o wiele większych kwot i, co za tym idzie, o wiele większych kosztów.

W dalszej kolejności zostanie podjęta próba uwzględnienia psucia się produktów dla przypadku utraconej sprzedaży, mieszaniny zaległych zamówień i utraconej sprzedaży dla pełnej informacji o kosztach oraz w modelach z ograniczeniami poziomu obsługi. Ostatecznie zostanie podjęta próba zastosowania tego podejścia w modelach gospodarowania zapasami z dokładnym ich wprowadzeniem, czyli bez przybliżeń.

Literatura

1. Dmytrów K. (2001). Propozycja deterministycznego modelu gospodarowania psującymi się zapasami. Zeszyty Naukowe. Uniwersytet Szczeciński, Szczecin, 320, 421-427.
2. Dmytrów K., Koczkodaj R. (2003). Zastosowanie modelu $\langle Q, r \rangle$ z mieszaniną zaległych zamówień i utraconej sprzedaży z ograniczeniami poziomu obsługi dla wyznaczenia optymalnego poziomu gotówki w bankomacie banku komercyjnego. W: Modelowanie preferencji a ryzyko '03. Red. T. Trzaskalik. AE, Katowice, 81-93.
3. Goyal S.K., Giri B.C. (2001). Recent Trends in Modeling of Deteriorating Inventory. European Journal of Operational Research, 134, 1-16.
4. Hadley G., Whitin T.M. (1963). Analysis of Inventory Systems. Prentice-Hall, Englewood Cliffs, NJ.

Monika Dyduch

Akademia Ekonomiczna w Katowicach

WSPÓŁCZYNNIKI TRANSFORMATY FALKOWEJ JAKO NARZĘDZIE GENERUJĄCE PROGNOZĘ PRZEDZIAŁOWĄ SZEREGÓW CZASOWYCH

Zasób narzędzi służących do analizy szeregów jest obecnie bardzo bogaty. Jednym z popularnych i szeroko znanych jest analiza Fouriera i analiza falkowa. Transformatą Fouriera umożliwia przeniesienie sygnału z dziedziny czasu do dziedziny częstotliwości. Jednakże przejście z układu czas-wartość do układu częstotliwość-wartość powoduje utratę informacji o czasie, tzn. nie można powiedzieć, kiedy dane zdarzenie częstotliwościowe miało miejsce. Natomiast transformata falkowa pozwala na przeniesienie sygnału z układu czas-wartość do układu czas-skala (czas-częstotliwość) dzięki czemu umożliwia analizę zmiany częstotliwości sygnału w funkcji czasu.

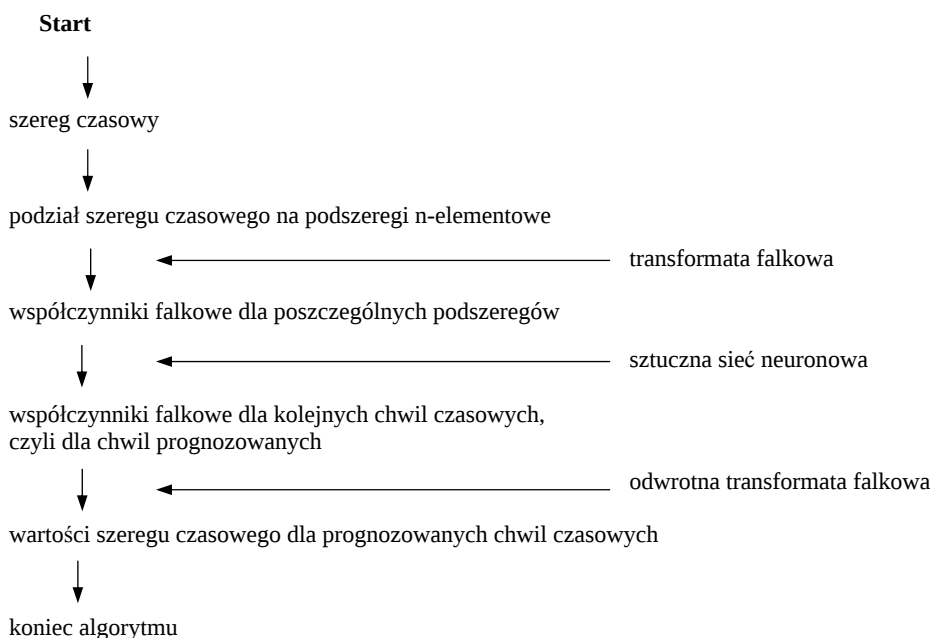
W opracowaniu integrujemy sieci neuronowe i transformatę falkową w celu dokonania predykcji szeregu czasowego. Dyskretną transformatę falkową przeprowadzamy przez tzw. algorytm „à trous”. Predykcję szeregu przeprowadzamy przez odwrotną transformatę falkową dla 8-elementowych przedziałów czasowych na podstawie współczynników transformaty falkowej wygenerowanych przez sieć neuronową.

1. Architektura modelu

W proponowanym modelu, którego architekturę prezentuje rys. 1, w pierwszej kolejności rzeczywisty szereg czasowy dzielimy na podszeregi n -elementowe, gdzie n jest wielokrotnością liczby 2. Z otrzymanego zbioru podszeregów do dalszych obliczeń wybieramy m początkowych próbek n -elementowych. Wyodrębnione m podszeregi poddajemy transformacie falkowej algorytmem „à Trous” z filtrem $h_3\left(\frac{1}{16}, \frac{1}{4}, \frac{3}{8}, \frac{1}{4}, \frac{1}{16}\right)$.

Kolejnym etapem jest inicjalizacja sieci neuronowej. Wykorzystujemy jedną z podstawowych własności sieci – zdolność do uogólniania wiedzy, czyli sieć nauczona na jednym zbiorze danych generuje właściwe wyniki dla innego zbioru danych nieuczestniczącego w procesie uczenia. Zatem poprzez sieć generujemy współczynniki falkowe przyszłych wartości szeregu, przyjmując za zbiór uczący współczynniki falkowe wcześniejszych obserwacji szeregu.

Z otrzymanych współczynników poprzez odwrotną transformatę falkową konstruujemy przyszłe (nowe) wartości szeregu czasowego.



Rys. 1. Architektura modelu

2. Algorytm „a trous”

Zakładamy, że dane $\{c_0(k)\}$ są iloczynem skalarnym z pikseli k funkcji $f(x)$ i funkcji skalującej, który odpowiada filtrowi dolnoprzepustowemu. Różnica sygnałów $\{c_0(k)\} - \{c_1(k)\}$ zawiera informacje pomiędzy dwoma skalami i dyskretnym zbiorem związanym z funkcją skalującą $\phi(x)$. Zatem spokrewnioną falką jest $\psi(x)$.

$$\frac{1}{2}\psi\left(\frac{x}{2}\right) = \phi(x) - \frac{1}{2}\phi\left(\frac{x}{2}\right) \quad (1)$$

Odległość pomiędzy próbkami powiększana przez czynnik 2 z skali $i-1$ do następnej i dyskretna transformata falkowa $w_i(k)$ wyrażamy przez

$$w_i(k) = c_{i-1}(k) - c_i(k)$$

gdzie

$$c_i(k) = \sum_l h(l) c_{i-1}(k + 2^{i-1}l)$$

Współczynniki $\{h(k)\}$ wyprowadzamy z funkcji skalującej $\phi(x)$

$$\frac{1}{2}\phi\left(\frac{x}{2}\right) = \sum_l h(l) \phi(x-l)$$

Algorytm pozwalający, by ponownie zbudować ramę danych jest oczywisty: ostatnia wygładzona macierz $\mathbf{c}_{n_p}$ jest dodawana do wszystkich $\mathbf{w}_i$.

$$c_0(k) = c_{n_p}(k) \sum_{j=1}^{n_p} w_j(k)$$

Wybierając liniową funkcję skalującą, czyli funkcję postaci

$$\phi(x) = \begin{cases} 1-|x| & \text{dla } x \in [-1,1] \\ 0 & \text{dla } x \notin [-1,1] \end{cases}$$

mamy

$$\frac{1}{2}\phi\left(\frac{x}{2}\right) = \frac{1}{4}\phi(x+1) + \frac{1}{2}\phi(x) + \frac{1}{4}\phi(x-1)$$

Wówczas $\mathbf{c}_1$ otrzymujemy z następującego wzoru

$$c_1(k) = \frac{1}{4}c_0(k-1) + \frac{1}{2}c_0(k) + \frac{1}{4}c_0(k+1)$$

Natomiast $\mathbf{c}_{j+1}$ przy wykorzystaniu $\mathbf{c}_j$ z

$$c_{j+1}(k) = \frac{1}{4}c_j(k-2^j) + \frac{1}{2}c_j(k) + \frac{1}{4}c_j(k+2^j)$$

Współczynniki falki z poziomu j są

$$c_{j+1}(k) = -\frac{1}{4}c_j(k-2^j) + \frac{1}{2}c_j(k) - \frac{1}{4}c_j(k+2^j)$$

Przedstawiony algorytm jest łatwo rozciągliwy do przestrzeni dwuwymiarowej. Prowadzi to do splotu z maską 3×3 pikseli dla falki spójnej z liniową interpolacją. Współczynniki maski są następujące:

$$\begin{pmatrix} \frac{1}{16} & \frac{1}{8} & \frac{1}{16} \\ \frac{1}{16} & \frac{1}{8} & \frac{1}{16} \\ \frac{1}{16} & \frac{1}{8} & \frac{1}{16} \end{pmatrix}$$

Dla każdego j otrzymujemy zbiór $\{w_j(k, l)\}$ który zawiera te same wartości pikseli, jak wyżej.

Jeżeli wybierzemy przestrzeń B_3 dla funkcji skalującej, to współczynniki splotu maski są postaci:

- w jednym wymiarze: $\left(\frac{1}{16}, \frac{1}{4}, \frac{3}{8}, \frac{1}{4}, \frac{1}{16}\right)$,
- w dwóch wymiarach:

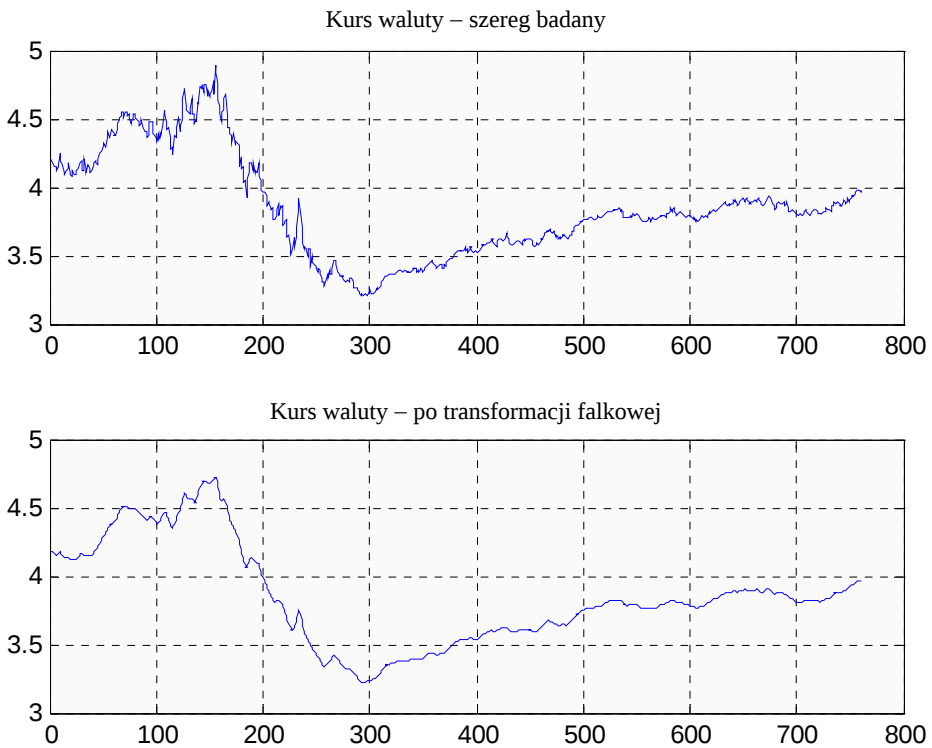
$$\begin{pmatrix} \frac{1}{256} & \frac{1}{64} & \frac{3}{128} & \frac{1}{64} & \frac{1}{256} \\ \frac{1}{256} & \frac{1}{64} & \frac{3}{128} & \frac{1}{64} & \frac{1}{256} \\ \frac{64}{3} & \frac{16}{3} & \frac{32}{9} & \frac{16}{3} & \frac{64}{3} \\ \frac{128}{1} & \frac{32}{1} & \frac{64}{3} & \frac{32}{1} & \frac{128}{1} \\ \frac{64}{1} & \frac{16}{1} & \frac{32}{3} & \frac{16}{1} & \frac{64}{1} \\ \frac{1}{256} & \frac{1}{64} & \frac{3}{128} & \frac{1}{64} & \frac{1}{256} \end{pmatrix}$$

3. Opis badania

W opracowaniu wykonujemy badania na szeregu czasowym prezentującym kurs wymiany euro w okresie 25.09.2006-25.09.2009 (wykres 1). Szereg zawiera 761 obserwacji. Wszelkie symulacje komputerowe i obliczenia wykonano w programie MATLAB opierając się na własnych autorskich programach.

Wykres 1

Kurs wymiany euro w okresie 25.09.2006-25.09.2009



Z uwagi na ograniczenia transformaty falkowej szereg poddany przekształceniu ograniczamy do 512 obserwacji, czyli do wielokrotności liczby 2.

Zgodnie z zaprezentowanym w pierwszym rozdziale algorytmem szereg czasowy składający się z 512 obserwacji dzielimy na podszeregi 8-elementowe (2^3), tzn. na podszeregi o elementach:

Podszereg 1: $t = 1, 2, 3, 4, 5, 6, 7, 8$

Podszereg 2: $t = 2, 3, 4, 5, 6, 7, 8, 9$

Podszereg 3: $t = 3, 4, 5, 6, 7, 8, 9, 10$

... ..

Podszereg 505: $t = 505, 506, 507, 508, 509, 510, 511, 512$

W wyniku podziału otrzymujemy 505 podszeregów 8-elementowych. Z otrzymanego zbioru podszeregów do dalszych obliczeń wybieramy 500 początkowych próbek 8-elementowych, natomiast 5 ostatnich pozostawiamy celem dokonania sprawdzenia działania algorytmu.

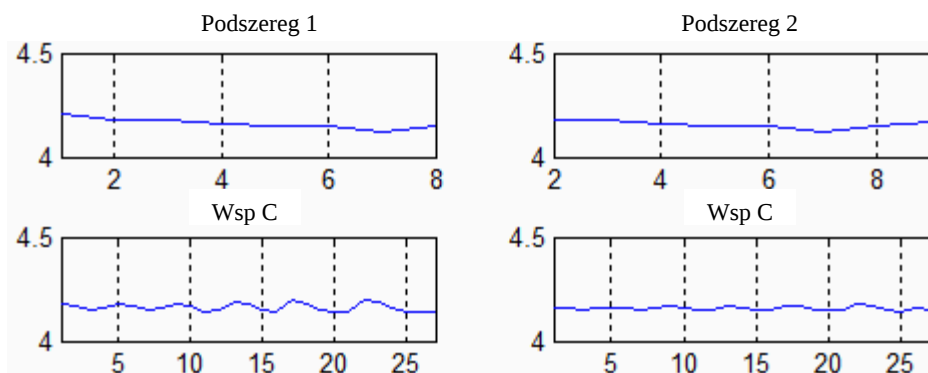
2.1. Transformaty falkowa

Każdy z 500 wyodrębnionych podszeregów 8-elementowych poddajemy transformacie falkowej algorytmem „a trous”, która pozwala przedstawić sygnał w postaci liniowej kombinacji współczynników $c_p(t), d_p(t)$. Zatem używając algorytmu „a trous” przeprowadzamy w programie komputerowym MATLAB pięciopoziomową dekompozycję każdego 8-elementowego podszeregu przyjmując zgodnie z rozważaniami zawartymi w rozdziale drugim filtr $h_3\left(\frac{1}{16}, \frac{1}{4}, \frac{3}{8}, \frac{1}{4}, \frac{1}{16}\right)$.

Dla podszeregu pierwszego i drugiego w wyniku transformaty falkowej algorytmem „a trous”, otrzymujemy odpowiednio współczynniki C pokazane na wykresie 2.

Wykres 2

Współczynniki C uzyskane po transformacie falkowej podszeregu 1 i 2

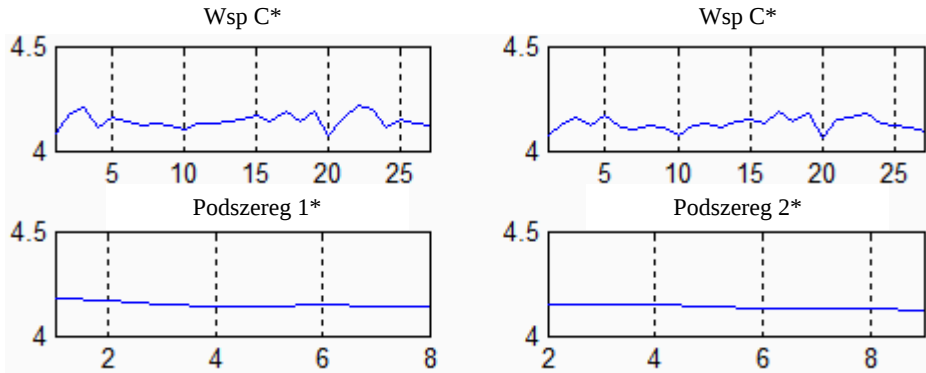


Analogicznie otrzymujemy współczynniki dla wszystkich pięciuset podszeregów 8-elementowych.

W wyniku odwrotnej transformaty falkowej dla podszeregu 1 i 2 otrzymujemy szeregi przedstawione na wykresie 3.

Wykres 3

Podszeregi 1 i 2 uzyskane poprzez odwrotną transformatę falkową



Otrzymane współczynniki są niezbędne do kolejnego etapu algorytmu z wykorzystaniem sztucznej sieci neuronowej.

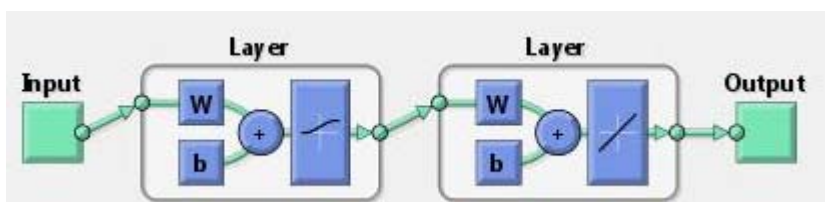
2.2. Sieć neuronowa

Głównym narzędziem prezentowanego algorytmu jest sieć neuronowa, która służy w przedstawianym algorytmie do generowania współczynników falkowych. Wygenerowane przez sieć współczynniki falkowe służą do konstrukcji przyszłych (nowych) wartości szeregu czasowego, które otrzymujemy poprzez odwrotną transformatę falkową.

Celem uruchomienia sieci przyjmujemy:

- wejście – 8-elementowe przedziały czasowe; w przypadku rozważanego szeregu będzie to przedział interesującej nas prognozy, czyli $t = 501, 502, 503, 504, 505, 506, 507, 508$ i $t = 502, 503, 504, 505, 506, 507, 508, 509$,
- zbiór uczący i testowy sieci; sieć uczy się na podszeregach 8-elementowych, uczymy sieć, że dla $t = 1, 2, 3, 4, 5, 6, 7, 8$ (podszereg 1) współczynniki falkowe przyjmują wartości pokazane na wykresie 2, uzyskane przez transformatę falkową, zatem sieć otrzymując na wejściu $t = 1, 2, 3, 4, 5, 6, 7, 8$ powinna wygenerować na wyjściu współczynniki falkowe zaprezentowane dla tego podszeregu na wykresie 2 itd.; dla pozostałych podszeregów sieć uczy się na 500 podszeregach,
- wyjście – współczynniki transformaty falkowej dla zadanej na wejściu próbki czasowej,
- parametry uczenia sieci:

- liczba epok = 1500,
 krok uczenia = 0.01
 dopuszczalny błąd sieci = $1e-12$,
- struktura sieci – strukturę sieci użytej w algorytmie przedstawia rys. 2.

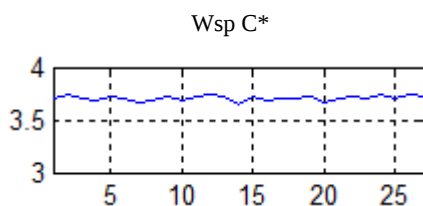


Rys. 2. Struktura sieci neuronowej użytej w algorytmie

Po uruchomieniu sieci przy wyżej określonych parametrach, sieć na wyjściu generuje współczynniki transformaty falkowej w formie macierzy o wymiarach 500×27 . Dla próbki 501 i 502 otrzymane współczynniki C prezentują wykresy 4 i 5.

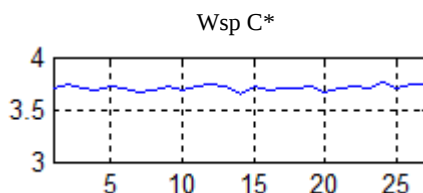
Wykres 4

Współczynniki falkowe wygenerowane przez sieć dla podszeregu 501



Wykres 5

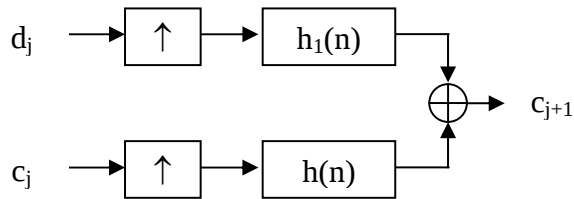
Współczynniki falkowe wygenerowane przez sieć dla podszeregu 502



2.3. Odwrotna transformata falkowa

Odwrotna transformata falkowa, której schemat blokowy przedstawia rys. 3, polega na kompozycji współczynników c_{j+1} na podstawie c_j i d_j

$$c_{j+1}(k) = \sum_m c_j(m)h(k-2m) + \sum_m d_j(m)h_1(k-2m)$$

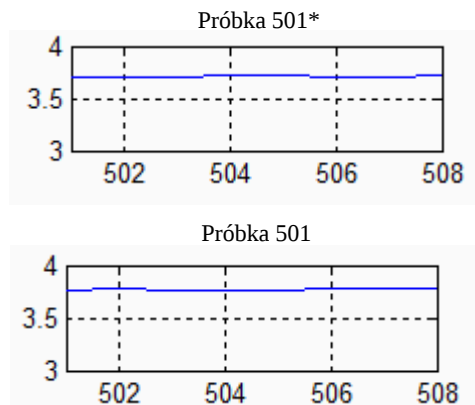


Rys. 3. Schemat blokowy odwrotnej dyskretnej transformaty falkowej

Zatem zgodnie z przedstawionym algorytmem, wykorzystując współczynniki falkowe wygenerowane przez sieć neuronową konstruujemy, poprzez odwrotną transformatę falkową, oryginalny szereg czasowy. Otrzymane podszeregi 501 i 502 przedstawiają wykresy 6 i 7.

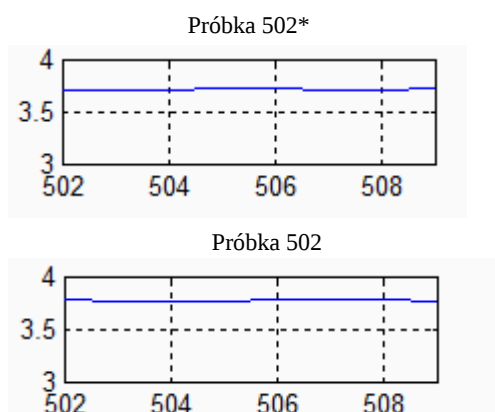
Wykres 6

Szereg czasowy otrzymany z odwrotnej transformaty falkowej
– próbka 501* oraz oryginalny szereg czasowy – próbka 501



Wykres 7

Szereg czasowy otrzymany z odwrotnej transformaty falkowej
– próbka 502^{*} oraz oryginalny szereg czasowy – próbka 502



Porównując na wykresie otrzymane wartości podszeregów z wartościami rzeczywistymi można stwierdzić, że wygenerowany przez algorytm szereg jest dobrym odzwierciedleniem szeregu rzeczywistego. Zatem przedstawiony algorytm jest skutecznym narzędziem w prognozowaniu szeregów czasowych.

Zakończenie

W pracy przedstawiono model bazujący na analizie falkowej i sztucznej sieci neuronowej. Skoncentrowano się wokół przedstawienia idei algorytmu prognozowania szeregów czasowych między innymi poprzez analizę falkową. Zaproponowano model prognozujący elementy szeregu czasowego na podstawie współczynników falkowych, generowanych przez sieć neuronową. Wyniki (wykresy 6 i 7) pokazują, że obserwacje wygenerowane przez zaproponowany model wiarygodnie odzwierciedlają rzeczywisty szereg czasowy, czyli model jest skutecznym narzędziem w prognozowaniu.

Literatura

1. Aussum et al (1997). Combing Neural Network Forecast on Wavelet-transformed Series [J]. Connection Science.
2. Burrus C.S., R Gopinath.A., Guo H. (1998). Introduction to Wavelets and Wavelet Transform. Prentice Hall.

-
3. Chui Ch.K. (1997). Wavelets: A Mathematical Tool for Signal Analysis. SIAM.
 4. Coakley J. and Fuertes A.-M. (2001). Nonparametric Cointegration Analysis of Real Exchange Rate [J]. Applied Financial Economics.
 5. Katsuragi H. (2000). Evidence of Multi-affinity in the Japanese Stock Market. [J]. Physica A.
 6. Misiti M., Misiti Y., Oppenheim G., Poggi J.-M. Wavelet Toolbox User's Guide Version 1.

Dorota Gawrońska

Politechnika Śląska

SZACOWANIE EFEKTYWNOŚCI PRZEDSIĘWZIĘCIA INWESTYCYJNEGO PRZEDSIĘBIORSTWA NA PODSTAWIE WSKAŹNIKA *NPVR* O ROZMYTYCH PARAMETRACH

Wstęp

Przed wyborem przedsięwzięcia inwestycyjnego do realizacji, należy dokonać analizy finansowej. Ma ona za zadanie wskazanie najbardziej efektywnego przedsięwzięcia inwestycyjnego spośród rozpatrywanych. Problemem jest tutaj brak danych na temat wszystkich potrzebnych do podjęcia decyzji parametrów finansowych, np. przepływów pieniężnych czy stopy dyskontowej w kolejnych okresach trwania inwestycji. Do oceny efektywności planowanych inwestycji stosuje się dyskontowane wskaźniki efektywności, takie jak wewnętrzna stopa zwrotu *IRR*, okres zwrotu *PB*, zaktualizowaną wartość netto *NPV* czy wskaźnik zaktualizowanej wartości netto *NPVR* itp. Charakteryzują one zależności pomiędzy przepływami finansowymi, jakie ta inwestycja może przynieść.

Problemem przy wyborze przedsięwzięcia inwestycyjnego na podstawie wskaźników finansowych jest często brak informacji o rozkładzie w czasie przepływów finansowych, wartości stopy dyskontowej w kolejnych okresach trwania inwestycji czy też niezgodność wartości wskaźników określanych przez ekspertów. W klasycznym podejściu obliczenie parametrów finansowych inwestycji rzeczowych bazuje na wartościach ostrych i pewnych. Niedoskonałością powyższego wskaźnika finansowego jest bazowanie na stałej wartości stopy dyskontowej w kolejnych okresach trwania inwestycji oraz przyjęciu dyskretnych wartości przepływów finansowych. Na tej podstawie zaproponowano uwzględnienie niepełnej informacji co do wartości przepływów finansowych i stopy dyskontowej w kolejnych okresach przebiegu przedsięwzięcia inwestycyjnego.

Przegląd literatury umożliwia zapoznanie się z pracami dotyczącymi uwzględniania niepewnej informacji dotyczącej wskaźników finansowych inwestycji, gdzie uwzględnia się niepewność informacyjną dotyczącą wartości samych wskaźników finansowych, a nie ich parametrów składowych [7]. W niniejszej pracy interpretacja niepewnej informacji następuje już na poziomie parametrów wskaźników finansowych, jak przepływy pieniężne czy stopa dyskontowa. Na podstawie rozmytych wartości tych parametrów określone są następnie rozmyte wartości wskaźników finansowych.

Przy określaniu optymalnego rozwiązania powinno się brać pod uwagę wartości wskaźników finansowych inwestycji określonych przez ekspertów. Jednym z nich jest wskaźnik *NPVR*. Celem niniejszej pracy jest przedstawienie wyboru najbardziej opłacalnego projektu na podstawie wskaźnika *NPVR* z uwzględnieniem rozmytych wartości przepływów pieniężnych oraz stopy dyskontowej. Skłania to do odpowiedniego sformułowania algorytmu określającego łączną ocenę wskaźnika *NPVR* dla wszystkich ekspertów biorących udział w analizie.

1. Wskaźnik *NPVR*

Przedsięwzięcia inwestycyjne, różne z punktu widzenia nakładów kapitałowych, porównać można na podstawie wskaźnika zaktualizowanej wartości netto *NPVR* (Net Present Value Ratio) określającego relacje wartości zaktualizowanej netto *NPV* do wartości obecnej nakładu inwestycyjnego koniecznego do realizacji przedsięwzięcia *PVI* (Present Value Of The Investment) [14]

$$NPVR = \frac{NPV}{PVI} \quad (1)$$

gdzie

NPV – wartość zaktualizowana (bieżąca) netto określona w następujący sposób

$$NPV = \sum_{t=1}^T \frac{CIF_t - COF_t}{(1+d)^t} \quad (2)$$

PVI – ujemne strumienie gotówki, dyskontowane na dany moment czasowy, liczony według formuły

$$PVI = \sum_{t=1}^T \frac{COF_t}{(1+d)^t} \quad (3)$$

Wartość *d* określa stopę dyskontową (stałą w kolejnych okresach trwania inwestycji), Natomiast *T* to czas trwania inwestycji.

Ze względu na niepewność informacyjną dotyczącą wartości przepływów pieniężnych w kolejnych okresach inwestycji, zwłaszcza przy założeniu długiego horyzontu realizacji inwestycji, powinno przyjąć się zmienne uwzględniające tę niepewność. Taką możliwość daje nam opracowana przez amerykańskiego profesora Lofti Zadeha teoria zbiorów rozmytych, umożliwiającą wyrażenie tej niepewności. Dlatego w dalszej części pracy zakłada się przyjęcie zmiennych określających poszczególne parametry wskaźników finansowych jako zmienne rozmyte.

2. Określenie wartości wskaźnika *NPVR* o rozmytych parametrach

W założeniu przyjmuje się rozmytość zmiennych dotyczących przepływów finansowych oraz stopy dyskontowej, z możliwością uwzględnienia zmiennej stopy dyskontowej w kolejnych okresach trwania inwestycji. Zakłada się więc, że

$$NPVR_{ij} = \frac{NPV_{ij}}{PVI_{ij}} \quad (4)$$

gdzie

$$PVI_{ij} = \sum_{t=1}^T \frac{COF_{ijt}}{\prod_{t_z=1}^t (1 + D_{ijt_z})} \quad (5)$$

$$NPV_{ij} = \sum_{t=1}^T \frac{CIF_{ijt} - COF_{ijt}}{\prod_{t_z=1}^t (1 + D_{ijt_z})} \quad (6)$$

CIF_{ijt} – strumień dodatnie gotówki (CIF – Cash Inflow) i -tego przedsięwzięcia określone przez j -tego eksperta dla okresu t ,

COF_{ijt} – ujemne strumienie gotówki (nakłady inwestycyjne) (COF – Cash Outflow) i -tego przedsięwzięcia określone przez j -tego eksperta dla okresu t .

W dalszej analizie zakłada się, że przy określaniu ocen bierze udział Q ekspertów, którzy oceniają N przedsięwzięć inwestycyjnych względem kryterium $NPVR$. Zadaniem jest znalezienie takiego przedsięwzięcia inwestycyjnego, dla którego osiągnięte zostanie maksimum oceny kryterium $NPVR$. Przyjęto określony okres trwania inwestycji T . Określony zostaje też zbiór badanych przedsięwzięć P

$$P = \{P_1, P_2, \dots, P_N\} \quad i = 1, \dots, N \quad (7)$$

oraz zbiór ekspertów E

$$E = \{E_1, E_2, \dots, E_Q\} \quad j = 1, \dots, Q \quad (8)$$

Do określenia wskaźnika $NPVR$ niezbędne są takie parametry, jak przepływy pieniężne CIF , COF oraz stopa dyskontowa D . Niżej przedstawione są założenia, co do tych parametrów.

Eksperci określają dodatnie i ujemne strumienie pieniężne CIF oraz COF . Przyjmuje się liczby rozmyte określające dodatnie strumienie pieniężne (CIF_{ijt}) oraz ujemne strumienie pieniężne (COF_{ijt}).

Do reprezentacji liczb rozmytych określających przepływy finansowe oraz stopę dyskontową zastosowano trójparametryczną reprezentację typu LR . Reprezentacja ta została zaproponowana przez D. Dubois i H. Prade [4]. Upraszcza ona znacznie wykonywanie operacji na liczbach rozmytych. Według J. Kacprzyka [6], liczba rozmyta $A \subseteq R$ jest liczbą rozmytą typu L - R wtedy i tylko wtedy, gdy

$$\mu_A(x) = \begin{cases} L\left(\frac{m-x}{\alpha}\right) & \text{dla } x < m \\ 1 & \text{dla } x = m \\ R\left(\frac{x-m}{\beta}\right) & \text{dla } x > m \end{cases} \quad (9)$$

Parametr m jest liczbą rzeczywistą zwaną wartością średnią ($\mu_A(m) = 1$), a α, β są odpowiednio „rozrzutem” lewostronnym i prawostronnym (left and right spreads), a L i R to funkcje odniesienia (reference function, shape function).

Reprezentacja LR jest charakteryzowana przez trzy parametry m, α, β , co zapisuje się jako $A = (m, \alpha, \beta)$. W artykule przyjęto następującą formułę określającą funkcje L oraz R [12]

$$L(x) = R(x) = \begin{cases} 0 & \text{dla } x < m - \alpha \\ 1 - |x|^p & \text{dla } m - \alpha \leq x \leq m + \beta \\ 0 & \text{dla } x > m + \beta \end{cases} \quad p > 0 \quad (10)$$

analogicznie dla R . W niniejszym opracowaniu przyjmuje się funkcje L oraz R określone powyższym wzorem, ze względu na fakt określania ocen w formie przedziałów wyrażających niepewność. Parametr p określa sposób zmiany wartości liczby w przedziałach $[m - \alpha, m]$ oraz $[m, m + \beta]$ (dla liniowej zmiany: parametr $p = 1$, dla nieliniowej zmiany: $p \neq 1$).

Operacje na liczbach rozmytych typu LR określane są jako operacje na tych trzech parametrach. Niżej przedstawione zostały wzory niezbędne do wykonania podstawowych obliczeń na liczbach rozmytych typu LR . Część podanych niżej wzorów będzie ścisła („=”), a część przybliżona („ $\cong$ ”) [11].

1. Jeżeli liczby rozmyte A_1 i A_2 przedstawimy w postaci trójek $A_1 = (m_{A_1}, \alpha_{A_1}, \beta_{A_1})$ i $A_2 = (m_{A_2}, \alpha_{A_2}, \beta_{A_2})$, to ich sumę w formie trójki

$$(A_1 + A_2) = (m_{A_1} + m_{A_2}, \alpha_{A_1} + \alpha_{A_2}, \beta_{A_1} + \beta_{A_2}) \quad (11)$$

Między parametrami sumy $(A_1 + A_2)$ i jej składników A_1 i A_2 zachodzą następujące zależności

$$m_{A_1+A_2} = m_{A_1} + m_{A_2} \quad (12)$$

$$\alpha_{A_1+A_2} = \alpha_{A_1} + \alpha_{A_2} \quad (13)$$

$$\beta_{A_1+A_2} = \beta_{A_1} + \beta_{A_2} \quad (14)$$

2. Jeżeli liczby rozmyte A_1 i A_2 przedstawimy w postaci trójek $A_1 = (m_{A_1}, \alpha_{A_1}, \beta_{A_1})$ i $A_2 = (m_{A_2}, \alpha_{A_2}, \beta_{A_2})$, to ich różnicę w formie trójki

$$(A_1 - A_2) = (m_{A_1} - m_{A_2}, \alpha_{A_1} + \beta_{A_2}, \alpha_{A_2} + \beta_{A_1}) \quad (15)$$

3. Iloczyn liczby rozmytej A_1 i A_2 : $(A_1 \cdot A_2)$ przedstawimy następująco

$$(A_1 \cdot A_2) = (m_{A_1} \cdot m_{A_2}, m_{A_1} \cdot \alpha_{A_2} + m_{A_2} \cdot \alpha_{A_1} - \alpha_{A_1} \cdot \alpha_{A_2}, m_{A_1} \cdot \beta_{A_2} + m_{A_2} \cdot \beta_{A_1} + \beta_{A_1} \cdot \beta_{A_2}) \quad (16)$$

4. Iloraz dwóch liczb rozmytych A_1 i A_2 , takich, że $A_1 = (m_{A_1}, \alpha_{A_1}, \beta_{A_1})$ i $A_2 = (m_{A_2}, \alpha_{A_2}, \beta_{A_2})$ oraz $A_1 > 0$ i $A_2 > 0$ przedstawimy następująco

$$(A_1 / A_2) \cong (m_{A_1}, \alpha_{A_1}, \beta_{A_1}) / (m_{A_2}, \alpha_{A_2}, \beta_{A_2}) = \left(\frac{m_{A_1}}{m_{A_2}}, \frac{m_{A_1} \cdot \beta_{A_2} + m_{A_2} \cdot \alpha_{A_1}}{m_{A_2} \cdot (m_{A_2} + \beta_{A_2})}, \frac{m_{A_1} \cdot \alpha_{A_2} + m_{A_2} \cdot \beta_{A_1}}{m_{A_2} \cdot (m_{A_2} - \alpha_{A_2})} \right) \quad (17)$$

gdzie

$$m_{A_1/A_2} = \frac{m_{A_1}}{m_{A_2}} \quad (18)$$

$$\alpha_{A_1/A_2} = \frac{m_{A_1} \beta_{A_2} + m_{A_2} \alpha_{A_1}}{m_{A_2} (m_{A_2} + \beta_{A_2})}, \quad m_{A_2} \neq 0 \quad (19)$$

$$\beta_{A_1/A_2} = \frac{m_{A_1} \alpha_{A_2} + m_{A_2} \beta_{A_1}}{m_{A_2} (m_{A_2} - \alpha_{A_2})}, \quad m_{A_2} - \alpha_{A_2} \neq 0, \quad m_{A_2} \neq 0 \quad (20)$$

5. Iloraz dwóch liczb rozmytych A_1 i A_2 takich, że $A_1 = (m_{A_1}, \alpha_{A_1}, \beta_{A_1})$ i $A_2 = (m_{A_2}, \alpha_{A_2}, \beta_{A_2})$ oraz $A_1 < 0$ i $A_2 > 0$ przedstawimy następująco

$$(A_1 / A_2) \cong (m_{A1}, \alpha_{A1}, \beta_{A1}) / (m_{A2}, \alpha_{A2}, \beta_{A2}) = \left(\frac{m_{A1}}{m_{A2}}, \frac{-m_{A1} \cdot \alpha_{A2} + m_{A2} \cdot \alpha_{A1}}{m_{A2} \cdot (m_{A2} - \alpha_{A2})}, \frac{-m_{A1} \cdot \beta_{A2} + m_{A2} \cdot \beta_{A1}}{m_{A2} \cdot (m_{A2} + \beta_{A2})} \right) \quad (21)$$

gdzie

$$\alpha_{A1/A2} = \frac{-m_{A1} \cdot \alpha_{A2} + m_{A2} \cdot \alpha_{A1}}{m_{A2} \cdot (m_{A2} - \alpha_{A2})} \quad (23)$$

$$\beta_{A1/A2} = \frac{-m_{A1} \cdot \beta_{A2} + m_{A2} \cdot \beta_{A1}}{m_{A2} \cdot (m_{A2} + \beta_{A2})} \quad (24)$$

Ocena dodatnich strumieni pieniężnych CIF_{ijt} i -tego przedsięwzięcia przez j -tego eksperta w czasie t jest modelowana za pomocą liczby rozmytej typu LR o następującej funkcji przynależności (porównaj ze wzorem (9))

$$\mu_{CIF_{ijt}}(cif_{ijt}) = \begin{cases} L\left(\frac{m_{CIF_{ijt}} - cif_{ijt}}{\alpha_{CIF_{ijt}}}\right) & \text{dla } cif_{ijt} < m_{CIF_{ijt}} \\ 1 & \text{dla } cif_{ijt} = m_{CIF_{ijt}} \\ R\left(\frac{cif_{ijt} - m_{CIF_{ijt}}}{\beta_{CIF_{ijt}}}\right) & \text{dla } cif_{ijt} > m_{CIF_{ijt}} \end{cases} \quad (25)$$

gdzie CIF_{ijt} określone jest charakterystyczną trójką $(m_{CIF_{ijt}}, \alpha_{CIF_{ijt}}, \beta_{CIF_{ijt}})$ oraz gdzie $\alpha_{CIF_{ijt}}, \beta_{CIF_{ijt}} > 0$ to ustalone rozrzuty lewo- i prawostronne (przedział określony przez eksperta $[cif_{ijt}^{\min}, cif_{ijt}^{\max}]$ wyrażający jego niepewność), $m_{CIF_{ijt}}$ to wartość ustalona przez eksperta jako najbardziej prawdopodobna bądź wartość liczona ze wzoru (26) w przypadku nie podania tej wartości przez eksperta, zaś L i R to ustalone funkcje bazowe (27) (porównaj ze wzorem (10)).

$$cif_{ijt}^{\text{mod}} = \frac{cif_{ijt}^{\min} + cif_{ijt}^{\max}}{2} \quad (26)$$

$$L(cif_{ijt}) = R(cif_{ijt}) = \begin{cases} 0 & \text{dla } cif_{ijt} < m_{CIF_{ijt}} - \alpha_{CIF_{ijt}} \\ 1 - |cif_{ijt}|^p & \text{dla } m_{CIF_{ijt}} + \beta_{CIF_{ijt}} \geq cif_{ijt} \geq m_{CIF_{ijt}} - \alpha_{CIF_{ijt}}, p > 0 \\ 0 & \text{dla } cif_{ijt} > m_{CIF_{ijt}} + \beta_{CIF_{ijt}} \end{cases} \quad (27)$$

Biorąc pod uwagę powyższe założenia można przyjąć, że

$$cif_{ijt}^{\min} = m_{CIF_{ijt}} - \alpha_{CIF_{ijt}} \quad (28)$$

$$cif_{ijt}^{\text{mod}} = m_{CIF_{ijt}} \text{ dla symetrycznej funkcji przynależności} \quad (29)$$

$$cif_{ijt}^{\text{max}} = m_{CIF_{ijt}} + \beta_{CIF_{ijt}} \quad (30)$$

Ocena ujemnych strumieni pieniężnych COF_{ijt} i -tego przedsięwzięcia przez j -tego eksperta w czasie t jest modelowana za pomocą liczby rozmytej typu LR o następującej funkcji przynależności (porównaj ze wzorem (9))

$$\mu_{COF_{ijt}}(cof_{ijt}) = \begin{cases} L\left(\frac{m_{COF_{ijt}} - cof_{ijt}}{\alpha_{COF_{ijt}}}\right) & \text{dla } cof_{ijt} < m_{COF_{ijt}} \\ 1 & \text{dla } cof_{ijt} = m_{COF_{ijt}} \\ R\left(\frac{cof_{ijt} - m_{COF_{ijt}}}{\beta_{COF_{ijt}}}\right) & \text{dla } cof_{ijt} > m_{COF_{ijt}} \end{cases} \quad (31)$$

gdzie COF_{ijt} określone jest charakterystyczną trójką $(m_{COF_{ijt}}, \alpha_{COF_{ijt}}, \beta_{COF_{ijt}})$ oraz gdzie $\alpha_{COF_{ijt}}, \beta_{COF_{ijt}} > 0$ to ustalone rozrzuty lewo- i prawostronne (przedział określony przez eksperta $[cof_{ijt}^{\min}, cof_{ijt}^{\max}]$, wyrażający jego niepewność), $m_{COF_{ijt}}$ to wartość ustalona przez eksperta jako najbardziej prawdopodobna bądź w przypadku braku jej podania liczona ze wzoru (32), zaś L i R to ustalone funkcje bazowe (33) (porównaj ze wzorem (10)).

$$cof_{ijt}^{\text{mod}} = \frac{cof_{ijt}^{\min} + cof_{ijt}^{\max}}{2} \quad (32)$$

$$L(cof_{ijt}) = R(cof_{ijt}) = \begin{cases} 0 & \text{dla } cof_{ijt} < m_{COF_{ijt}} - \alpha_{COF_{ijt}} \\ 1 - |cof_{ijt}|^p & \text{dla } m_{COF_{ijt}} + \beta_{COF_{ijt}} \geq cof_{ijt} \geq m_{COF_{ijt}} - \alpha_{COF_{ijt}}, \quad p > 0 \\ 0 & \text{dla } cof_{ijt} > m_{COF_{ijt}} + \beta_{COF_{ijt}} \end{cases} \quad (33)$$

Biorąc pod uwagę powyższe założenia można przyjąć, że

$$cof_{ijt}^{\min} = m_{COF_{ijt}} - \alpha_{COF_{ijt}} \quad (33)$$

$$cof_{ijt}^{\text{mod}} = m_{COF_{ijt}} \text{ dla symetrycznej funkcji przynależności} \quad (34)$$

$$cof_{ijt}^{\max} = m_{COF_{ijt}} + \beta_{COF_{ijt}} \quad (35)$$

Ocena kosztu kapitału D_{ijt_z} i -tego przedsięwzięcia j -tego eksperta w czasie t_z jest modelowana za pomocą liczby rozmytej typu LR o następującej funkcji przynależności (porównaj ze wzorem (9))

$$\mu_{D_{ijt_z}}(d_{ijt}) = \begin{cases} L\left(\frac{m_{D_{ijt_z}} - d_{ijt}}{\alpha_{D_{ijt_z}}}\right) & \text{dla } d_{ijt} < m_{D_{ijt_z}} \\ 1 & \text{dla } d_{ijt} = m_{D_{ijt_z}} \\ R\left(\frac{d_{ijt} - m_{D_{ijt_z}}}{\beta_{D_{ijt_z}}}\right) & \text{dla } d_{ijt} > m_{D_{ijt_z}} \end{cases} \quad (36)$$

gdzie D_{ijt_z} określone jest charakterystyczną trójką $(m_{D_{ijt_z}}, \alpha_{D_{ijt_z}}, \beta_{D_{ijt_z}})$ oraz gdzie $\alpha_{D_{ijt_z}}, \beta_{D_{ijt_z}} > 0$ to ustalone rozrzuty lewo- i prawostronne (określone przez eksperta jako wyraz niepewności $[d_{ijt}^{\min}, d_{ijt}^{\max}]$), $m_{D_{ijt_z}}$ to wartość ustalona przez eksperta jako najbardziej prawdopodobna, bądź liczona ze wzoru (37), zaś L i R to ustalone funkcje bazowe (38)).

$$d_{ijt}^{\text{mod}} = \frac{d_{ijt}^{\min} + d_{ijt}^{\max}}{2} \quad (37)$$

$$L(d_{ijt}) = R(d_{ijt}) = \begin{cases} 0 & \text{dla } d_{ijt} < m_{D_{ijt_z}} - \alpha_{D_{ijt_z}} \\ 1 - |d_{ijt}|^p & \text{dla } m_{D_{ijt_z}} + \beta_{D_{ijt_z}} \geq d_{ijt} \geq m_{D_{ijt_z}} - \alpha_{D_{ijt_z}}, p > 0 \\ 0 & \text{dla } d_{ijt} > m_{D_{ijt_z}} + \beta_{D_{ijt_z}} \end{cases} \quad (38)$$

Zważywszy na przyjęte założenia odnośnie do przedziału wartości podanej przez eksperta otrzymujemy

$$d_{ijt}^{\min} = m_{D_{ijt_z}} - \alpha_{D_{ijt_z}} \quad (39)$$

$$d_{ijt}^{\text{mod}} = m_{D_{ijt_z}} \text{ dla symetrycznej funkcji przynależności} \quad (40)$$

$$d_{ijt}^{\max} = m_{D_{ijt_z}} + \beta_{D_{ijt_z}} \quad (41)$$

Mając określone wielkości strumieni pieniężnych czy ogólnie przepływów pieniężnych dla każdego eksperta wyznacza się oceny rozmyte przedsięwzięć inwestycyjnych $NPV_{ij}(m_{NPV_{ij}}, \alpha_{NPV_{ij}}, \beta_{NPV_{ij}})$ zgodnie ze wzorem

$$NPV_{ij} = \sum_{t=1}^T \frac{CIF_{ijt} - COF_{ijt}}{\prod_{t_z=1}^t (1 + D_{ijt_z})} \quad (42)$$

W zależności od określonej funkcji L oraz R dla przepływów finansowych COF_{ijt} , CIF_{ijt} i D_{ijt} , które tworzą liczbę rozmytą NPV_{ij} funkcję przynależności oceny NPV_{ij} określamy następująco:

1. Dla $p = 1$ [11].

Do określenia funkcji przynależności wartości wskaźnika NPV_{ij} należy określić parametry $(m_{NPV_{ij}}, \alpha_{NPV_{ij}}, \beta_{NPV_{ij}})$. Zakładając wartość parametru $p = 1$ dla funkcji L oraz R można przyjąć w przybliżeniu trójkątną funkcję przynależności $\mu_{NPV_{ij}}(npv_{ij})$

$$\mu_{NPV_{ij}}(npv_{ij}) = \begin{cases} L\left(\frac{m_{NPV_{ij}} - npv_{ij}}{\alpha_{NPV_{ij}}}\right) & \text{dla } npv_{ij} < m_{NPV_{ij}} \\ 1 & \text{dla } npv_{ij} = m_{NPV_{ij}} \\ R\left(\frac{npv_{ij} - m_{NPV_{ij}}}{\beta_{NPV_{ij}}}\right) & \text{dla } npv_{ij} > m_{NPV_{ij}} \end{cases} \quad (43)$$

Funkcje L i R można wtedy zapisać w następujący sposób

$$L(npv_{ij}) = R(npv_{ij}) = \begin{cases} 0 & \text{dla } npv_{ij} < m_{NPV_{ij}} - \alpha_{NPV_{ij}} \\ 1 - |npv_{ij}| & \text{dla } m_{NPV_{ij}} + \beta_{NPV_{ij}} \geq npv_{ij} \geq m_{NPV_{ij}} - \alpha_{NPV_{ij}} \\ 0 & \text{dla } npv_{ij} > m_{NPV_{ij}} + \beta_{NPV_{ij}} \end{cases} \quad (44)$$

Stosowanie opracowanych wzorów dla operacji arytmetycznych na liczbach rozmytych typu LR znacznie upraszcza nam rzeczywistą funkcję przynależności i dlatego, aby dokładnie określić funkcję przynależności, należy skorzystać z poziomej reprezentacji liczb rozmytych, polegającą na przedstawieniu zbioru rozmytego za pomocą tzw. α – przekrojów tego zbioru. Zbiorem poziomu α (α – level set) lub α – przekrojem zbioru A (α – cut) nazywamy zbiór określony przez funkcję charakterystyczną

$$\chi_{A_\alpha} = \begin{cases} 1 & \text{dla } \mu_A(x) \geq \alpha \\ 0 & \text{dla } \mu_A(x) < \alpha \end{cases} \quad (45)$$

Jeśli przez αA_α oznaczymy zbiór rozmyty o funkcji przynależności $\mu(x) = \alpha \wedge \chi_{A_\alpha}$, zbiór rozmyty A możemy wyrazić jako sumę zbiorów rozmytych αA_α po wszystkich poziomach α

$$A = \bigcup_{\alpha \in \Lambda_A} \alpha A_\alpha \quad (46)$$

Wszystkie operacje na liczbach rozmytych A i B mogą zostać sprowadzone do operacji na zbiorach ostrych przedziałów odpowiadających α -przekrojom [13]

$$(A \otimes B)_\alpha = A_\alpha \otimes B_\alpha \quad (47)$$

Z tego powodu operacje na liczbach rozmytych mogą zostać zdefiniowane w następującej postaci ($\min$ oraz $\max$ oznaczają granice danego przedziału na danym poziomie α) [2]

a) suma dwóch liczb rozmytych A i B

$$\begin{aligned} c_{\min} &= a_{\min} + b_{\min} \\ c_{\max} &= a_{\max} + b_{\max} \end{aligned} \quad (48)$$

b) różnica dwóch liczb rozmytych A i B

$$\begin{aligned} c_{\min} &= a_{\min} - b_{\max} \\ c_{\max} &= a_{\max} - b_{\min} \end{aligned} \quad (49)$$

c) iloczyn dwóch liczb rozmytych A i B

$$\begin{aligned} c_{\min} &= \min[a_{\min} \cdot b_{\min}, a_{\min} \cdot b_{\max}, a_{\max} \cdot b_{\min}, a_{\max} \cdot b_{\max}] \\ c_{\max} &= \max[a_{\min} \cdot b_{\min}, a_{\min} \cdot b_{\max}, a_{\max} \cdot b_{\min}, a_{\max} \cdot b_{\max}] \end{aligned} \quad (50)$$

d) iloraz dwóch liczb rozmytych A i B

$$\begin{aligned} c_{\min} &= \min[a_{\min} / b_{\min}, a_{\min} / b_{\max}, a_{\max} / b_{\min}, a_{\max} / b_{\max}] \\ c_{\max} &= \max[a_{\min} / b_{\min}, a_{\min} / b_{\max}, a_{\max} / b_{\min}, a_{\max} / b_{\max}] \end{aligned} \quad (51)$$

Pojęcie α -przekrojów znajduje zastosowanie ze względu na to, że ułatwia identyfikację funkcji przynależności. Na podstawie odpowiedniej ilości α -przekrojów można odtworzyć z żadaną dokładnością funkcje przynależności zbioru rozmytego. Znając przynależność elementów rozważań do poszczególnych α -przekrojów możemy określić przybliżoną funkcję przynależności zbioru rozmytego. Przeprowadzając obliczenia dla odpowiednio dużej liczby α -przekrojów według powyższych zasad uzyskuje się względnie dokładne przybliżenie postaci funkcji przynależności.

2. Dla $p > 0, p \neq 1$.

W celu przedstawienia wartości NPV_{ij} przy wykorzystaniu poziomej reprezentacji liczb rozmytych należy dokonać podziału liczb rozmytych COF_{ijt} , CIF_{ijt} i D_{ijt_z} na α -przekroje i przeprowadzić obliczenia na wartościach z granic przedziałów. Granice przedziału dla poziomu α można określić na podstawie wzorów

$$x_{L\alpha} = m - ((1 - \alpha_i) \cdot \alpha^p)^{1/p}, \quad \text{dla} \quad x_{L\alpha} \in [m - \alpha, m] \quad (52)$$

$$x_{P\alpha} = m + ((1 - \alpha_i) \cdot \beta^p)^{1/p}, \quad \text{dla} \quad x_{P\alpha} \in [m, m + \beta] \quad (53)$$

gdzie:

$x_{L\alpha}$ – lewa granica przedziału,

$x_{P\alpha}$ – prawa granica przedziału,

α_i – kolejne wartości α -poziomu,

p – parametr z funkcji L i R.

Po określeniu wartości liczby rozmytej NPV_{ij} , należy określić wartość PVI zgodnie ze wzorem (5) przy wykorzystaniu przedstawionych powyżej wzorów.

Dla każdego eksperta wyznacza się ocenę rozmytą wskaźnika $NPVR_{ij}$ i -tego przedsięwzięcia, modelowaną za pomocą liczby rozmytej $NPVR_{ij}$. Przy korzystaniu ze wzorów arytmetycznych charakterystycznych dla liczb rozmytych typu LR liczba $NPVR_{ij}$ opisana jest charakterystyczną trójką $m_{NPVR_{ij}}, \alpha_{NPVR_{ij}}, \beta_{NPVR_{ij}}$, gdzie $\alpha_{NPVR_{ij}}, \beta_{NPVR_{ij}} > 0$ to ustalone rozrzuty lewo- i prawostronne, $m_{NPVR_{ij}}$ to wartość średnia. Funkcje L oraz R opisane są następująco

$$L(npvr_{ij}) = R(npvr_{ij}) = \begin{cases} 0 & \text{dla} \quad npvr_{ij} < m_{NPVR_{ij}} - \alpha_{NPVR_{ij}} \\ 1 - |npvr_{ij}| & \text{dla} \quad m_{NPVR_{ij}} + \beta_{NPVR_{ij}} \geq npvr_{ij} \geq m_{NPVR_{ij}} - \alpha_{NPVR_{ij}} \\ 0 & \text{dla} \quad npvr_{ij} > m_{NPVR_{ij}} + \beta_{NPVR_{ij}} \end{cases} \quad (54)$$

Określając funkcję przynależności oceny $NPVR_{ij}$ można przyjąć w przybliżeniu, że będzie to trójkątna funkcja przynależności $\mu_{NPVR_{ij}}(npvr_{ij})$

$$\mu_{NPVR_{ij}}(npvr_{ij}) = \begin{cases} L\left(\frac{m_{NPVR_{ij}} - npvr_{ij}}{\alpha_{NPVR_{ij}}}\right) & \text{dla} \quad npvr_{ij} < m_{NPVR_{ij}} \\ 1 & \text{dla} \quad npvr_{ij} = m_{NPVR_{ij}} \\ R\left(\frac{npvr_{ij} - m_{NPVR_{ij}}}{\beta_{NPVR_{ij}}}\right) & \text{dla} \quad npvr_{ij} > m_{NPVR_{ij}} \end{cases} \quad (55)$$

Takie podejście nie gwarantuje nam idealnego rozwiązania. Dlatego dla dokładnego określenia funkcji przynależności należy skorzystać z poziomej reprezentacji liczb rozmytych.

Dla $p \neq 1, p > 0$ chcąc określić funkcję przynależności wartości wskaźnika $NPVR$ bazujemy na poziomej reprezentacji liczb rozmytych NPV oraz PVI . Wartości α -przekrojów odpowiadają wartościom funkcji przynależności $\mu_{NPVR_{ij}}(npvr_{ij})$.

Mając określone wartości $NPVR$ dla każdego eksperta obliczamy ocenę łączną dla każdego przedsięwzięcia i zgodnie ze wzorem

$$NPVR_i = \frac{\sum_{j=1}^Q NPVR_{ij}}{Q} \quad (56)$$

przy założeniu jednakowego zaufania do ekspertów.

Poszukując maksymalnej wartości rozmytej kryterium $NPVR$ dla każdego przedsięwzięcia ze zbioru N należy dokonać jej defuzyfikacji (wyostrzeniu) stosując jedną z metod defuzyfikacji [11]. Uzyskujemy wtedy rzeczywistą wartość oceny przedsięwzięcia względem kryterium $NPVR$: $NPVR_i(P_i)$. W niniejszym opracowaniu zaprezentowano metodę środka ciężkości jako metodę defuzyfikacji.

Przy wyborze przedsięwzięcia inwestycyjnego bierze się pod uwagę wiele wskaźników finansowych, głównie NPV czy IRR . Analizując wartości wskaźnika $NPVR$ rozpatrywanych przedsięwzięć inwestycyjnych za korzystniejszą uznaje się inwestycję o największej wartości wskaźnika $NPVR$ [8]. Optymalizacja sprowadza się wtedy do zadania poszukiwania rozwiązania, dla którego wartość oceny kryterium $NPVR$ jest maksymalna

$$NPVR_i(P_i) \rightarrow MAX \quad (57)$$

Najkorzystniejszym przedsięwzięciem będzie to, które osiągnie największą wartość $NPVR$.

Przy wyborze optymalnego przedsięwzięcia inwestycyjnego należy wziąć pod uwagę również inne wskaźniki określające efektywność inwestycji.

Niżej przedstawiony został przykład analizy dwóch przedsięwzięć inwestycyjnych o następujących przepływach pieniężnych: dodatnich CIF i ujemnych COF oraz stopie dyskontowej D . Przyjmuje się, że w analizie bierze udział dwóch ekspertów, zmiana wartości liczby w przedziałach $[m - \alpha, m]$ oraz $[m, m + \beta]$ jest liniowa ($p = 1$) oraz nie określono wartości najbardziej prawdopodobnej m .

Tabela 1

Wartości przepływów finansowych i stopy dyskontowej przedsięwzięcia inwestycyjnego

Eksperti	Okres	Parametr	Wartości ocen	
			Min	Max
Ekspert 1	1	CIF	0	0
		COF	30000	40000
		D	0	0,02
	2	CIF	5000	8000
		COF	0	0
		D	0,01	0,02
	3	CIF	40000	80000
		COF	0	0
		D	0,02	0,03
	4	CIF	30000	80000
		COF	0	0
		D	0,04	0,05
	5	CIF	20000	40000
		COF	0	0
		D	0,04	0,05
Ekspert 2	1	CIF	0	0
		COF	30000	45000
		D	0	0,01
	2	CIF	5000	8000
		COF	0	0
		D	0,01	0,02
	3	CIF	35000	65000
		COF	0	0
		D	0,02	0,04
	4	CIF	40000	75000
		COF	0	0
		D	0,05	0,07
	5	CIF	30000	70000
		COF	0	0
		D	0,06	0,09

Przyjmując poziomą reprezentację zbioru rozmytego otrzymujemy następujące wartości liczb rozmytych na poszczególnych poziomach (do analizy przyjęto 10 poziomów).

W tabeli 2 przedstawione są wartości na poszczególnych poziomach przepływów finansowych i stopy dyskontowej dla pierwszego okresu analizy. W analogiczny sposób określa się wartości liczb rozmytych parametrów finansowych dla kolejnych okresów. Na podstawie określonych przepływów finansowych i stopy dyskontowej określa się wartości kolejnych czynników wskaźnika NPV, których wartości przedstawione są w tabeli 3.

Tabela 2

Granice przekrojów przepływów pieniężnych i czynnika dyskontowego

Wartość α -poziomu	Wartość lewej strony przedziału wartości oceny CIF	Wartość prawej strony przedziału wartości oceny CIF	Wartość lewej strony przedziału wartości oceny COF	Wartość prawej strony przedziału wartości oceny COF	Wartość lewej strony przedziału wartości czynnika dyskontowego	Wartość prawej strony przedziału wartości czynnika dyskontowego	Wartość lewej strony wartości różnicy CIF-COF	Wartość prawej strony wartości różnicy CIF-COF
1	0	0	35 000,00	35 000,00	1,010	1,010	-35 000,00	-35 000,00
0,9	0	0	34 500,00	35 500,00	1,009	1,011	-35 500,00	-34 500,00
0,8	0	0	34 000,00	36 000,00	1,008	1,012	-36 000,00	-34 000,00
0,7	0	0	33 500,00	36 500,00	1,007	1,013	-36 500,00	-33 500,00
0,6	0	0	33 000,00	37 000,00	1,006	1,014	-37 000,00	-33 000,00
0,5	0	0	32 500,00	37 500,00	1,005	1,015	-37 500,00	-32 500,00
0,4	0	0	32 000,00	38 000,00	1,004	1,016	-38 000,00	-32 000,00
0,3	0	0	31 500,00	38 500,00	1,003	1,017	-38 500,00	-31 500,00
0,2	0	0	31 000,00	39 000,00	1,002	1,018	-39 000,00	-31 000,00
0,1	0	0	30 500,00	39 500,00	1,001	1,019	-39 500,00	-30 500,00
0	0	0	30 000,00	40 000,00	1,000	1,020	-40 000,00	-30 000,00

Tabela 3

Wartości czynników wskaźnika NPV na poszczególnych poziomach dla obu ekspertów

Ekspert	Poziom	Okres1		Okres1		Okres2		Okres2		Okres3		Okres3		Okres4		Okres4		Okres5	
		MIN	MAX	MIN	MAX	MIN	MAX	MIN	MAX	MIN	MAX	MIN	MAX	MIN	MAX	MIN	MAX	MIN	MAX
Ekspert 1	1	-34653,50	-34653,50	6340,54	6340,54	6340,54	6340,54	57100,51	57100,51	57100,51	57100,51	50088,16	50088,16	50088,16	50088,16	23965,63	23965,63	23965,63	23965,63
	0,9	-34835,00	-33786,80	6233,74	6447,64	6233,74	6447,64	55088,54	59120,28	55088,54	59120,28	47694,52	52493,34	47694,52	52493,34	22809,43	22809,43	25128,45	25128,45
	0,8	-35395,70	-33231,30	6127,26	6555,06	6127,26	6555,06	53084,34	61147,90	53084,34	61147,90	45312,37	54910,11	45312,37	54910,11	21659,83	21659,83	26297,95	26297,95
	0,7	-35958,60	-32678,00	6021,09	6662,78	6021,09	6662,78	51087,89	63183,38	51087,89	63183,38	42941,64	57338,52	42941,64	57338,52	20516,79	20516,79	27474,14	27474,14
	0,6	-36523,70	-32126,70	5915,22	6770,83	5915,22	6770,83	49099,15	65226,77	49099,15	65226,77	40582,28	59778,63	40582,28	59778,63	19380,27	19380,27	28657,06	28657,06
	0,5	-37090,90	-31577,60	5809,66	6879,18	5809,66	6879,18	47118,09	67278,09	47118,09	67278,09	38234,25	62230,50	38234,25	62230,50	18250,24	18250,24	29846,76	29846,76
	0,4	-37660,30	-31030,60	5704,41	6987,86	5704,41	6987,86	45144,69	69337,37	45144,69	69337,37	35897,49	64694,18	35897,49	64694,18	17126,67	17126,67	31043,27	31043,27
Ekspert 2	1	-37313,40	-37313,40	6372,08	6372,08	6372,08	6372,08	47588,35	47588,35	47588,35	47588,35	51628,87	51628,87	51628,87	51628,87	41762,49	41762,49	41762,49	41762,49
	0,9	-37889,20	-36367,10	6267,84	6476,52	6267,84	6476,52	46070,30	49112,25	46070,30	49112,25	51364,79	54910,53	51364,79	54910,53	42894,66	42894,66	46833,35	46833,35
	0,8	-38670,60	-35589,50	6163,81	6581,17	6163,81	6581,17	44558,08	50642,03	44558,08	50642,03	49606,97	56698,55	49606,97	56698,55	40947,47	40947,47	48825,06	48825,06
	0,7	-39453,50	-34813,50	6059,97	6686,02	6059,97	6686,02	43051,66	52177,70	43051,66	52177,70	47859,10	58496,75	47859,10	58496,75	39014,90	39014,90	50831,81	50831,81
	0,6	-40238,00	-34038,90	5956,34	6791,08	5956,34	6791,08	41551,02	53719,29	41551,02	53719,29	46121,12	60305,18	46121,12	60305,18	37096,87	37096,87	52853,72	52853,72
	0,5	-41024,10	-33265,90	5852,91	6896,34	5852,91	6896,34	40056,14	55266,83	40056,14	55266,83	44933,00	62123,90	44933,00	62123,90	35193,26	35193,26	54890,89	54890,89
	0,4	-41811,60	-32494,40	5749,68	7001,87	5749,68	7001,87	38566,99	56820,34	38566,99	56820,34	42674,66	63952,98	42674,66	63952,98	33303,97	33303,97	56943,43	56943,43
	0,3	-42600,80	-31724,40	5646,65	7107,48	5646,65	7107,48	37083,56	58379,84	37083,56	58379,84	40966,05	65792,46	40966,05	65792,46	31428,91	31428,91	59011,44	59011,44
	0,2	-43391,50	-30955,90	5543,84	7213,36	5543,84	7213,36	35605,81	59945,36	35605,81	59945,36	39267,13	67642,40	39267,13	67642,40	29567,97	29567,97	61095,05	61095,05
	0,1	-44183,70	-30188,90	5441,20	7319,45	5441,20	7319,45	34133,73	61516,92	34133,73	61516,92	37577,84	69502,87	37577,84	69502,87	27721,06	27721,06	63194,35	63194,35
	0	-44977,50	-29423,40	4853,43	7920,79	4853,43	7920,79	32667,29	63094,54	32667,29	63094,54	34891,64	69334,66	34891,64	69334,66	24008,01	24008,01	61049,39	61049,39

Stosując wzór (48) otrzymujemy rozmyte wartości wskaźnika *NPV* (tabela 4).

Tabela 4

Wartości wskaźnika *NPV* na poszczególnych poziomach

Eksperti	Wartość α -poziomu	Wartość lewej strony przedziału <i>NPVR</i>	Wartość prawej strony przedziału <i>NPVR</i>
Ekspert 1	1	102 841,37	102 841,37
	0,9	96 991,24	109 402,96
	0,8	90 788,08	115 679,71
	0,7	84 608,80	121 980,87
	0,6	78 453,27	128 306,56
	0,5	72 321,36	134 656,92
	0,4	66 212,95	141 032,08
	0,3	60 127,90	147 432,17
	0,2	54 066,09	153 857,34
	0,1	48 027,39	160 307,71
	0	41 531,09	167 278,48
Ekspert 2	1	110 038,36	110 038,36
	0,9	108 708,40	120 965,51
	0,8	102 605,72	127 157,26
	0,7	96 532,09	133 378,80
	0,6	90 487,32	139 630,34
	0,5	84 471,24	145 912,06
	0,4	78 483,66	152 224,16
	0,3	72 524,40	158 566,83
	0,2	66 593,28	164 940,27
	0,1	60 690,12	171 344,67
	0	51 442,85	171 975,94

Na podstawie wzoru (4) określone zostają wartości wskaźnika *NPVR* na poszczególnych poziomach (tabela 5).

Tabela 5

Wartości wskaźnika *NPVR* na poszczególnych poziomach
uzyskana dla obu ekspertów

Eksperti	Wartość α -poziomu	Wartość lewej strony przedziału <i>NPVR</i>	Wartość prawej strony przedziału <i>NPVR</i>
1	2	3	4
Ekspert 1	1	2,94	2,94
	0,9	2,73	3,17
	0,8	2,52	3,40
	0,7	2,32	3,64
	0,6	2,12	3,89
	0,5	1,93	4,14

cd. tabeli 5

1	2	3	4
	0,4	1,74	4,41
	0,3	1,56	4,68
	0,2	1,39	4,96
	0,1	1,22	5,26
	0	1,04	5,58
Ekspert 2	1	2,93	2,93
	0,9	2,84	3,29
	0,8	2,63	3,53
	0,7	2,43	3,78
	0,6	2,23	4,05
	0,5	2,05	4,32
	0,4	1,87	4,61
	0,3	1,70	4,92
	0,2	1,53	5,24
	0,1	1,37	5,57
	0	1,14	5,73

Zakładając jednakowe zaufanie do ekspertów otrzymujemy ostateczną postać wskaźnika NPVR w postaci liczby rozmytej (tabela 6).

Tabela 6

Wartość wskaźnika NPVR na poszczególnych poziomach

Wartość α -poziomu	Wartość lewej strony przedziału NPVR	Wartość prawej strony przedziału NPVR
1	2,94	2,94
0,9	2,79	3,23
0,8	2,58	3,47
0,7	2,37	3,71
0,6	2,18	3,97
0,5	1,99	4,23
0,4	1,81	4,51
0,3	1,63	4,80
0,2	1,46	5,10
0,1	1,29	5,41
0	1,09	5,65

Stosując metodę środka ciężkości jako metodę defuzyfikacji otrzymujemy rzeczywistą wartość wskaźnika NPVR: 3,14.

Podsumowanie

Analiza finansowa inwestycji rzeczowych jest zadaniem skomplikowanym [10]. Wymaga bowiem do jej określenia danych dotyczących przyszłości. Ponieważ inwestycje z reguły są długoterminowe trudno jest przewidzieć sytuację na rynku. Dlatego informacja, na podstawie której podejmuje się decyzje dotyczące inwestycji jest nieprecyzyjna. Ogranicza to znacznie skuteczność i efektywność różnych metod projektowania, sterowania, modelowania, prognozowania itd. Jednym z proponowanych rozwiązań może być teoria zbiorów rozmytych [3]. Dzięki przyjęciu rozmytości parametru kosztu kapitału oraz przepływów pieniężnych w kolejnych okresach trwania inwestycji można dokładniej określić opłacalność przedsięwzięcia. Aby jednak dokonać głębszej analizy efektywności badanych przedsięwzięć należy wziąć pod uwagę również inne wskaźniki opłacalności przedsięwzięć [14].

Literatura

1. Chojcan J. (2001). Zbiory rozmyte i ich zastosowanie. Politechnika Śląska, Gliwice.
2. Drewniak J. (2001). Liczby rozmyte. W: Zbiory rozmyte i ich zastosowanie. Red. J. Chojcan i J. Łęski. Politechnika Śląska, Gliwice.
3. Driankow D., Hellendoorn H., Reinfrank M. (1996). Wprowadzenie do sterowania rozmytego. WNT, Warszawa.
4. Dubois D., Prade H. (1980). Fuzzy Set and Systems – Theory and Applications. Academic Press, New York.
5. Kacprzyk J. (1986). Zbiory rozmyte w analizie systemowej. PWN, Warszawa.
6. Kacprzyk J. (2001). Wieloetapowe sterowanie rozmyte. WNT, Warszawa.
7. Kuchta D. (2001). Miękką matematyką w zarządzaniu – zastosowanie liczb przedziałowych i rozmytych w rachunkowości zarządczej. Politechnika Wrocławska, Wrocław.
8. Kurek W. (1997). Efektywność inwestycji rzeczowych w gospodarce rynkowej. UMCS, Rzeszów.
9. Łachwa A. (2001). Rozmyty świat zbiorów, liczb, relacji, faktów, reguł i decyzji. AOW Exit, Warszawa.
10. Marcinek K. (2000). Finansowa ocena przedsięwzięć inwestycyjnych przedsiębiorstw. AE, Katowice.
11. Piegat A. (1999). Modelowanie i sterowanie rozmyte. AOW Exit, Warszawa.

-
12. Rutkowski L. (2005). Metody i techniki sztucznej inteligencji. PWN, Warszawa.
 13. Walaszek-Babiszewska A. (2004). Statistical and Fuzzy Modelling of Grain Material Sampling and Operations. Politechnika Śląska, Gliwice.
 14. Wilczek M.T. (2004). Podstawy zarządzania projektem inwestycyjnym. AE, Katowice.

Paweł Hanczar

Uniwersytet Ekonomiczny we Wrocławiu

TARYFY STREFOWE W PROBLEMACH WYZNACZANIA TRAS POJAZDÓW

Wstęp

W opracowaniu rozważono możliwość zastosowania modeli optymalizacyjnych do wyznaczania planu dostaw w przypadku strefowego systemu rozliczania kosztów. W pierwszej części zaprezentowano wady i zalety strefowego systemu rozliczania oraz jego wykorzystanie w sieciach logistycznych. Następnie przedyskutowano możliwości wspomagania procesu planowania tras w takim systemie z wykorzystaniem metod optymalizacyjnych. Rozważone zostały dwa podejścia, pierwsze wykorzystujące model przydziału oraz drugie korzystające ze sformułowania modelu wyznaczania tras dostaw. Wprowadzone modele zostały użyte do rozwiązania zadania rzeczywistego. Pracę kończy opis integracji opracowanej metody z systemem informatycznym wykorzystywanym w badanym przedsiębiorstwie.

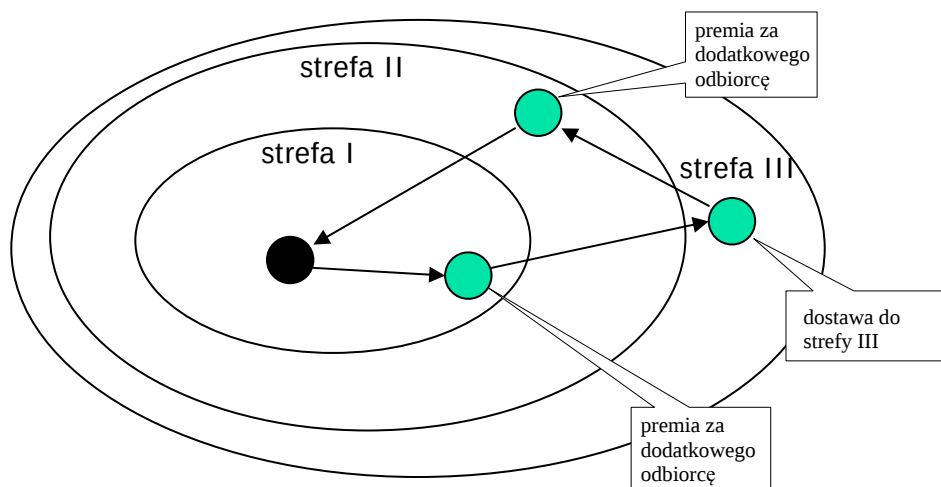
Strefowy system rozliczania kosztów dostaw jest stosowany w dystrybucji bardzo często. Główną przyczyną jego popularności jest duża łatwość użycia. Koszty transportu zależą tutaj od wagi ładunku oraz strefy, do której jest realizowana dostawa. Określenie kosztów dostawy według opisywanego systemu wymaga jedynie odczytania wielkości opłaty z odpowiednio przygotowanej tabeli. W przypadku ładunków typu LTL (less-than-truckload), kiedy to dostawca może w jednym kursie obsłużyć więcej niż jednego odbiorcę, koszty dostawy ponoszone przez odbiorcę pozostają bez zmian. W takich sytuacjach dostawca obniża koszty realizacji dostaw dzięki łączeniu w jednej trasie dostaw do więcej niż jednego odbiorcy.

Postępująca specjalizacja w zakresie realizacji usług dystrybucyjnych prowadzi do sytuacji, w których przedsiębiorstwa zlecają realizację dostaw zewnętrznym firmom transportowym. Zastosowanie w takich przypadkach taryfy strefowej jest bardziej korzystne dla firmy transportowej, gdyż, jak już wspomniano, odpowiednie planowanie tras dostaw umożliwia obsługę kilku odbiorców realizując jedną trasę, pomimo że każdy ładunek opłacany jest

według taryfy strefowej. Ze względu na fakt, że koszty transportu stanowią jedną z ważniejszych pozycji kosztowych firmy próbują wprowadzać indywidualne rozwiązania w tym obszarze. Przykładem mogą być tutaj strategie mieszane polegające na korzystaniu z usług firm transportowych rozliczających się według różnych typów taryf transportowych [4] czy stosowanie w rozliczeniach taryf bardziej korzystnych dla firmy zlecającej transport niż dla firmy transportowej.

1. Taryfa strefowa uwzględniająca liczbę odbiorców

Jedna z najprostszych modyfikacji standardowej taryfy strefowej polega na uwzględnieniu w rozliczeniach najdalszej obsługiwanej strefy oraz liczby odbiorców, do których jest dostarczany towar. Firma zlecająca transport płaci w takim przypadku standardową stawkę strefową wyłącznie za dostawę do najdalszego odbiorcy oraz stałą kwotę za każdego dodatkowo obsługiwanego klienta (rys. 1). Ponadto, ustala się jak bardzo oddalone mogą być lokalizacje obsługiwane w jednej trasie.



Rys. 1. Zmodyfikowana taryfa strefowa

Przedstawiony sposób rozliczeń zapewnia, analogicznie jak w przypadku klasycznej taryfy strefowej, dużą łatwość w wyznaczaniu kosztów dostaw ponoszonych przez odbiorców oraz możliwość bezpośredniego powiązania przychodów za zrealizowane dostawy z wydatkami dla firm transportowych. Jednak tak zdefiniowana taryfa wymaga skomplikowanego sposobu ustalania tras dostaw. Im lepszy plan dostaw, tym niższe koszty ich realizacji, co przy stałym przychodzie (każdy odbiorca opłaca transport zgodnie ze stawką bez względu na sposób realizacji dostawy) bezpośredni przekłada się na zysk przedsiębiorstwa.

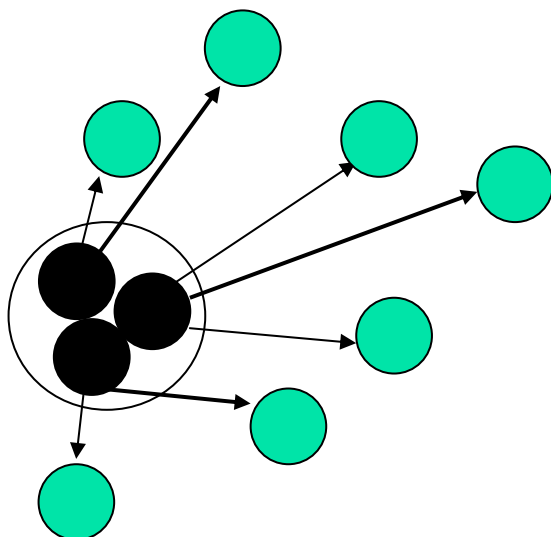
Przedstawione w dalszej części artykułu rozważania nie dotyczą wyłącznie takich układów jak zaprezentowane powyżej powiązanie dostawcy z firmą transportową i odbiorcami. Mają one zastosowanie także podczas planowania dostaw przez firmy korzystające z transportu zewnętrznego.

2. Sformułowanie liniowe problemu planowania dostaw

Opisywany problem jest zbliżony do dwóch zagadnień rozważanych bardzo gruntownie w literaturze. Pierwsze z nich to zagadnienie przydziału (general assignment problem), drugie to wyznaczanie tras dostaw (vehicle routing problem). Dostosowanie wymienionych zagadnień do prezentowanej sytuacji wymaga jednak kilku modyfikacji. W przypadku pierwszego sformułowania należy uwzględnić ładowność pojazdu oraz zmodyfikować funkcję celu. W drugim zagadnieniu jest niezbędna modyfikacja funkcji celu. Dodatkowo należy pamiętać, że wielkość rozwiązywanych zadań dla drugiego sformułowania jest mocno ograniczona – nie więcej niż około 100 lokalizacji. Ze względu na fakt, że na początkowym etapie badań trudno jednoznacznie stwierdzić, które podejście pozwoliłoby uzyskać najlepsze wyniki w dalszych badaniach rozważono równoległe wykorzystanie obu przedstawionych podejść.

2.1. Modyfikacja zagadnienia przydziału

Jako pierwsze zostanie rozważone sformułowanie bazujące na modyfikacji zagadnienia przydziału. Koncepcja użycia tego podejścia polega na przydzieleniu wszystkich odbiorców do tras (część klasyczna). Dodatkowo należy zapewnić, że w żadnej trasie nie będzie przekroczona ładowność pojazdu oraz zagwarantować, że wartość funkcji celu będzie obliczana zgodnie z warunkami podanymi w opisie problemu (rys. 2).



Rys. 2. Problem planowania dostaw w strefach jako zadanie przydziału

W modelu zaprezentowanym wzorami (1)-(6) wykorzystuje się dwie grupy zmiennych decyzyjnych x_{ij} oraz z_j . Pierwsza to zmienna binarna, określa czy odbiorca i jest obsługiwany w trasie j . Natomiast druga to koszt obsługi najdalszego odbiorcy w trasie j . Ponadto, przyjęto następujące oznaczenia parametrów: c_i to koszt dostawy do odbiorcy i , q_i to zapotrzebowanie odbiorcy i , d_{ik} to odległość lokalizacji i i k oraz $d_{\max}$ to maksymalna odległość pomiędzy odbiorcami w jednej trasie.

$$\min \sum_{j=1}^m z_j \quad (1)$$

$$z_j - c_i x_{ij} \geq 0 \quad \text{dla każdego } i, j \quad (2)$$

$$\sum_{i=1}^n q_i x_{ij} \leq Q \quad \text{dla każdego } j \quad (3)$$

$$\sum_{j=1}^m x_{ij} \geq 1 \quad \text{dla każdego } i \quad (4)$$

$$d_{ik} (x_{ij} + x_{kj} - 1) \leq d_{\max} \quad \text{dla każdego } i, k, j \quad (5)$$

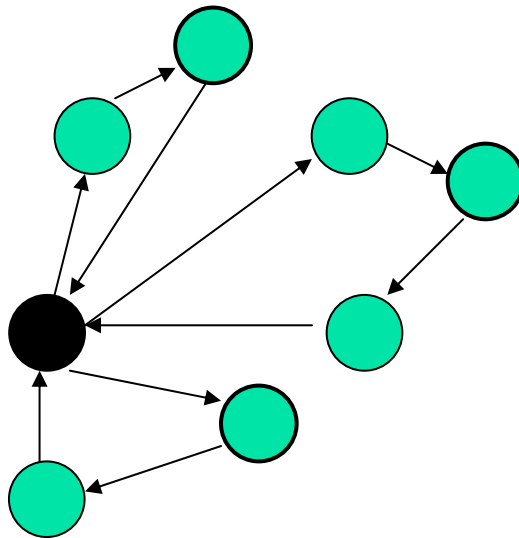
$$x_{ij} \in \{0, 1\} \quad \text{dla każdego } i, j \quad (6)$$

Funkcja celu (1) wraz z ograniczeniem (2) zapewnia, że koszt zrealizowania wszystkich dostaw będzie równy sumie kosztów dostaw do najdalszej lokalizacji w każdej trasie. Ograniczenie (3) zapewnia, że ładowność pojazdu na żadnej trasie nie zostanie przekroczona. Ograniczenie (4) gwarantuje, że odbiorca zostanie obsłużony dokładnie raz. Ograniczenie (5) zostało wprowadzone do modelu w celu zapewnienia, że odległość pomiędzy każdą parą odbiorców w jednej trasie nie przekroczy zadanej odległości maksymalnej. Ostatnie ograniczenie (6) gwarantuje, że zmienna decyzyjna x_{ij} będzie zmienną binarną.

Przedstawione sformułowanie zostało użyte do rozwiązania losowych zadań testowych. Największe rozwiązane zadanie składało się z 16 miast, a czas znajdowania rozwiązania wyniósł 550 sekund. Dodatkowo, przebieg procesu optymalizacji uniemożliwiał skrócenie czasu znajdowania rozwiązania, gdyż we wszystkich rozwiązywanych zadaniach zaobserwowano dużą różnicę pomiędzy przybliżeniem liniowym zadania całkowitoliczbowego. Nie podjęto próby rozwiązywania zadania rzeczywistego za pomocą zaprezentowanego sformułowania.

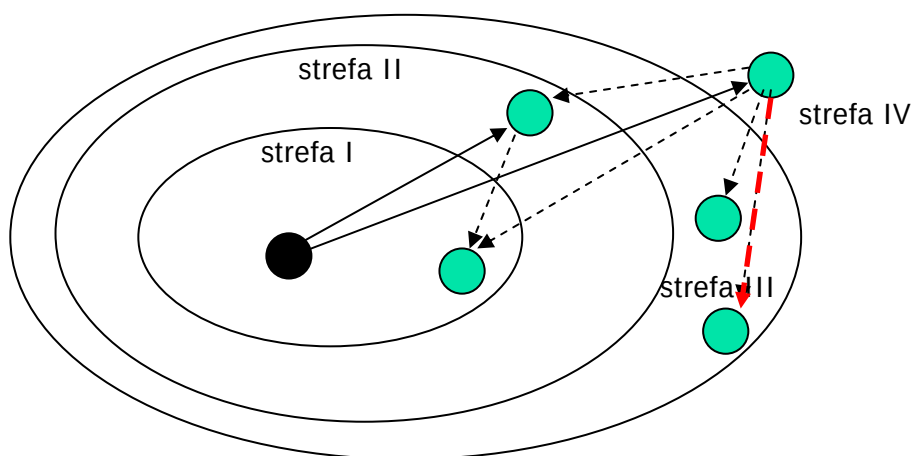
2.2. Modyfikacja zagadnienia wyznaczania tras dostaw

W drugiej kolejności zostanie rozważone sformułowanie bazujące na modyfikacji zagadnienia wyznaczania tras dostaw. Koncepcja użycia tego podejścia polega na wygenerowaniu tras dostaw (analogicznie jak w podstawowej postaci problemu wyznaczania tras dostaw), ale przy zmienionej postaci funkcji celu. Koszt obsługi trasy powinien wynosić tyle samo ile koszt obsługi strefy najdalszego odbiorcy w trasie (rys. 3).



Rys. 3. Problem planowania dostaw w strefach jako zadanie wyznaczania tras dostaw

W zadaniu wyznaczania tras dostaw wartość funkcji celu jest sumą długości połączeń użytych w rozwiązaniu. Powoduje to, że planowana poprzednio wyłącznie modyfikacja funkcji celu w modelu wyjściowym okazała się niemożliwa. Model zaprezentowany w dalszej części został opracowany od podstaw, a jedyną częścią wspólną ze sformułowaniem wyznaczania tras dostaw jest sposób zdefiniowania zmiennej decyzyjnej x_{ij} . Zmienna ta w wyjściowym sformułowaniu (model wyznaczania tras dostaw) określa czy odcinek (i,j) został użyty w rozwiązaniu optymalnym. Przyjmuje wartość 1, gdy pojazd korzysta z połączenia (i,j) w rozwiązaniu oraz 0 w pozostałych przypadkach. W nowym sformułowaniu interpretacja zmiennej jest inna, gdyż połączenia i ich wagi zostały wykorzystane w wyznaczaniu wartości funkcji celu. W nowym sformułowaniu rozważa się połączenia skierowane. Zostały wyróżnione dwa typy połączeń. Pierwszy z nich to połączenia od centrum dystrybucji do odbiorcy, którym przypisano wagę równą kosztowi obsługi strefy danego odbiorcy. Drugi typ to połączenia pomiędzy odbiorcami. Połączeniom tego typu przypisano wagę równą kosztowi obsługi dodatkowej lokalizacji w trasie. W lokalizacji każdego odbiorcy musi kończyć się dokładnie jedno połączenie co oznacza, że będzie obsłużony dokładnie raz. Ponadto, od odbiorców połączonych bezpośrednio z centrum dystrybucji mogą rozpoczynać się połączenia do innych odbiorców oddalonych od odbiorcy mniej niż $d_{\max}$ oraz należących do tej samej lub bliższej strefy niż dany odbiorca. Takie sformułowanie modelu gwarantuje, że koszt obsługi trasy będzie kosztem obsługi najdalej położonego odbiorcy w trasie (rys. 4).



Rys. 4. Sposób sformułowania modelu zapewniający możliwość wyznaczenia wartości funkcji celu

W modelu zaprezentowanym wzorami (7)-(11), jak wspomniano już wcześniej, wykorzystuje się zmienną decyzyjną x_{ij} . Przyjmuje ona wartość 1, gdy odcinek $\langle i,j \rangle$ jest użyty w rozwiązaniu oraz 0 w pozostałych przypadkach. Dodatkowo przyjęto następujące oznaczenia parametrów: c_i to koszt dostawy do odbiorcy i , q_i to zapotrzebowanie odbiorcy i oraz L to maksymalna liczba odbiorców obsługiwanych w jednej trasie. W sformułowaniu modelu bezpośrednio nie wykorzystuje się parametru $d_{\max}$ (jak w poprzednim podejściu), gdyż jest ono wykorzystywane podczas przygotowywania modelu. Jeśli odległość pomiędzy odbiorcami przekracza $d_{\max}$, to połączenie takie nie jest uwzględniane podczas procesu optymalizacji. Takie podejście powoduje zmniejszenie rozmiaru modelu, co może mieć duży wpływ na możliwości jego rozwiązywania.

$$\min \sum_{i=1}^n \sum_{j=1}^n c_{ij} x_{ij} \quad (7)$$

$$\sum_{i=0}^n x_{ij} = 1 \quad \text{dla każdego } j > 0 \quad (8)$$

$$\sum_{j=1}^n x_{ij} \leq L x_{0i} \quad \text{dla każdego } i > 0 \quad (9)$$

$$\sum_{j=1}^n q_j x_{ij} + q_i x_{0i} \leq Q \quad \text{dla każdego } i > 0 \quad (10)$$

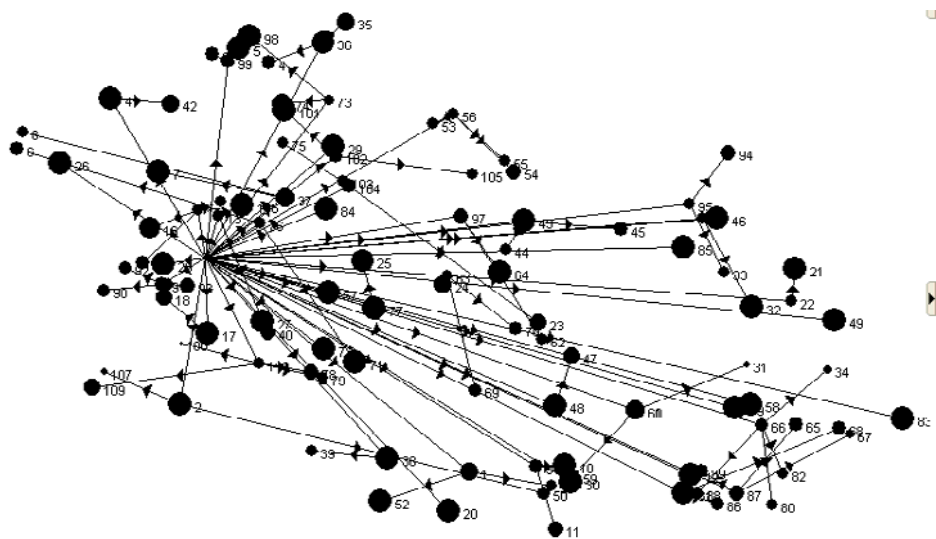
$$x_{ij} \in \{0,1\} \quad \text{dla każdego } i,j \quad (11)$$

Wartość funkcji celu (7) jest sumą wag odcinków użytych w rozwiązaniu. Zaprezentowany wcześniej sposób definiowania połączeń i ich wag gwarantuje, że wartość funkcji celu jest zgodna z warunkami zadania. Ograniczenie (8) zapewnia, że każdy odbiorca (z wyjątkiem centrum dystrybucji, które w modelu oznaczono indeksem 0) będzie obsługiwany dokładnie raz. Ograniczenia (9) powoduje, że połączenia rozpoczynające się od odbiorcy (interpretowane jako obsługa dodatkowej lokalizacji w trasie) są dopuszczane tylko od odbiorców obsługiwanych bezpośrednio z centrum dystrybucji (połączenie x_{0i}). Dodatkowo ograniczenie to zapewnia, że liczba obsługiwanych dodatkowo lokalizacji nie przekroczy L . Ograniczenie (10) gwarantuje, że ładowność pojazdu nie zostanie przekroczona na żadnej z tras. Ostatnie ograniczenie (11) zapewnia, że zmienna decyzyjna x_{ij} będzie zmienną binarną.

Przedstawione sformułowanie zostało przetestowane na losowo generowanych zadaniach. Model doskonale radzi sobie z rozwiązywaniem zadań składających się nawet z 100 miast. Czas rozwiązania tak dużych zadań wyniósł około 10 minut. Ponadto stwierdzono, że przebieg procesu optymalizacji umoż-

liwiało skrócenie czasu znajdowania rozwiązania. We wszystkich rozwiązywanych zadaniach zaobserwowano, że różnica pomiędzy przybliżeniem liniowym zadania całkowitoliczbowego już na początku rozwiązywania problemu jest bardzo mała. Ustawienie dokładności rozwiązania na poziomie 1% spowodowało, że zadania zawierające nawet 120 odbiorców rozwiązywane były w ciągu kilku sekund.

Następnie model ten został użyty do rozwiązania zadania rzeczywistego składającego się ze 110 miast. Rozwiązane zadania z zadaną dokładnością na poziomie 1% uzyskano w ciągu 2 sekund, co w porównaniu do jednego dnia pracy dwóch osób odpowiedzialnych za planowanie tras stanowi dużą oszczędność. Koszt realizacji dostaw w znalezionym rozwiązaniu jest o 2% mniejszy niż w rozwiązaniu proponowanym przez przedsiębiorstwo. Tak mała różnica jest najprawdopodobniej spowodowana subiektywnym podejściem do ograniczenia maksymalnej odległości odbiorców obsługiwanych w jednej trasie oraz możliwością przenoszenia do planowych tras odbiorców z przyszłych okresów (czyli wcześniejsza realizacja dostawy po uzgodnieniu z odbiorcą). W modelu założenie jest spełnione dla każdej trasy, w rzeczywistości osoby planujące trasy mają tutaj pewien możliwość wydłużania tego dystansu. Tak więc poluzowanie ograniczenia, o którym mowa, oraz możliwość przeniesienia do bieżącego okresu odbiorcy w celu zwiększenia załadowania pojazdu znacznie ułatwiają proces planowania dystrybucji. Dokładna ocena jakości proponowanych rozwiązań w stosunku do rozwiązań wyznaczanych przez dział logistyki wymaga przeanalizowania większej liczby planów dystrybucji. Graf wyznaczonego rozwiązania zaprezentowano na rys. 5.



Rys. 5. Graf rozwiązania rzeczywistego problemu strefowego planowania dostaw

Zakończenie

O możliwości zastosowania proponowanych rozwiązań optymalizacyjnych bardzo często decyduje możliwość oprogramowania i integracji algorytmu rozwiązywania z działającym w przedsiębiorstwie systemem informatycznym. W omawianym przypadku zaproponowany model został zaimplementowany w wersji wykorzystującej komercyjne oprogramowanie CPLEX oraz w wersji działającej z darmowym solverem COIN-MP. Na etapie implementacji przedsiębiorstwo nie podjęło jeszcze decyzji czy zdecyduje się na zakup oprogramowania komercyjnego.

Oprogramowanie zostało zintegrowane z systemem SAP R3 przy wykorzystaniu standardowych metod, takich jak tzw. punkty wyjścia do rozwiązań klienckich (user-exit) oraz zdalne wywołanie funkcji (remote function call) pozwalające na wywoływanie zewnętrznych programów. Taki sposób integracji zapewnił pełny komfort korzystania (użytkownik końcowy nie wie, że w chwili planowania dostaw uruchamiane są zewnętrzne procedury) oraz umożliwił wykorzystanie profesjonalnego oprogramowania optymalizacyjnego (CPLEX, COIN-MP).

Zaprezentowane badania pokazały ponadto, że w obszarze optymalizacji modele uwzględniające specyficzne warunki zadania pozwalają uzyskiwać rozwiązania bliskie optymalnemu dla zadań rzeczywistych. Stawiają natomiast pod znakiem zapytania rozwiązania wszelkiego typu podejścia uniwersalne.

Literatura

1. Bell W., Dalberto L., Fisher M., Greeneld A., Jaikumar R., Kedia P., Mack R., Prutzman P. (1983). Improving the Distribution of Industrial Gases with an On-line Computerized Routing and Scheduling Optimizer. *Interfaces*, 6, 1983.
2. Chien T., Balakrishnan A., Wong R. (1989). An Integrated Inventory Allocation and Vehicle Routing Problem. *Transportation Science*, 2.
3. Golden B., Assad A., Dahl R. (1984). Analysis of a Large Scale Vehicle Routing Problem with Inventory Component. *Large Scale Systems*, 7.
4. Hanczar P. (2008). Transport Planning in Conditions of Different Transport Tariffs – Application of Integer Programming. *Total Logistics Management*, 1. AGH Press, Kraków.

Katarzyna Jakowska-Suwalska

Politechnika Śląska

ZAGADNIENIE LOKALIZACJI OBIEKTÓW MODELOWANYCH JAKO SYSTEMY G/G/1 Z RÓŻNYMI KLASAMI KLIENTÓW

Wstęp

W pracy zajęto się zagadnieniem lokalizacji p obiektów traktowanych jako systemy masowej obsługi G/G/1 do obsługi n obiektów umieszczonych w znanej sieci. Każdy obsługiwany obiekt traktowany jest jako odrębne źródło zgłoszeń. W zagadnieniu lokalizacji rozróżnia się więc różne klasy klientów obsługiwanych przez systemy. Zagadnienie rozwiązano metodą podziału i ograniczeń gdzie za funkcję celu przyjęto minimalizację średniego całkowitego czasu obsługi w systemach. Do obliczeń wykorzystano aproksymację dyfuzyjną.

Zagadnienie lokalizacji dotyczy organizacji ruchu pomiędzy wieloma obiektami, gdzie jedno z nich to obiekty obsługiwane, pozostałe to obiekty obsługujące. Najczęściej problem dotyczy takiego rozmieszczenia obiektów obsługujących, aby znajdowały się one jak najbliżej obiektów obsługiwanych. W większości problemów zakłada się, że obiekty powinny znajdować się na znanej sieci połączeń (drogowych, telekomunikacyjnych itp.). W przypadku, gdy obiekty obsługujące to centra serwisowe (pogotowie ratunkowe, straż pożarna, bank krwi, pogotowie serwisowe), należy je zlokalizować w takich punktach sieci, aby czas dostępu do obiektów obsługiwanych był jak najmniejszy. W takich przypadkach centra serwisowe powinny być umieszczone jak najbliżej obiektów obsługiwanych, a dodatkowo czasy oczekiwania na usługę były jak najmniejsze. W tego typu zagadnieniach należy centra serwisowe modelować za pomocą systemów masowej obsługi.

Po raz pierwszy w zagadnieniu lokalizacji zastosowano teorię kolejek w pracy [22]. W pracach [7; 7; 17; 18] poszukiwano lokalizacji w sieci pojedynczego systemu M/G/1.

Dla różnych postaci funkcji celu stosowano różne heurystyki dla znalezienia optymalnej lokalizacji. Lokalizację systemu M/G/1 na znanej sieci drzew rozważano w pracach [4; 10].

W artykule [1] analizowano lokalizację systemu M/G/1 w przypadku, gdy system odrzuca pewne zgłoszenia.

W pracach [15; 18] analizowano lokalizację systemu M/G/k (z wieloma stanowiskami obsługi). W artykule [21] rozważano model optymalnej lokalizacji wielu systemów M/G/1 w znanej sieci połączeń.

Przegląd prac dotyczących lokalizacji stacji ambulansów (z trzydziestu ostatnich lat) zamieszczono w Luce Brotcorne, Gilbert Laporte and Frédéric Semet (2003). Modele lokalizacji podzielono na trzy grupy: modele deterministyczne, stochastyczne i dynamiczne.

Opisy i zastosowanie różnych heurystyk używanych w zagadnieniach lokalizacji można znaleźć w [11; 12; 14].

Analizę modelu lokalizacji system M/M/1/r przeprowadzono w pracach [15; 16].

2. System masowej obsługi G/G/1

W systemie masowej obsługi G/G/1 zgłoszenia przybywają do stanowiska obsługi w niezależnych od siebie odstępach czasu, których rozkład ma wartość oczekiwaną $\frac{1}{\lambda}$ oraz wariancję σ_A^2 . Czas obsługi zgłoszenia ma rozkład

z wartością oczekiwaną $\frac{1}{\mu}$ i wariancją σ_B^2 . Po obsłudze, zgłoszenia opuszczają

więc system w odstępach czasu, których rozkład ma wartość oczekiwaną $\frac{1}{\mu}$

i wariancję σ_B^2 . Na podstawie centralnego twierdzenia granicznego możemy przyjąć, że liczba klientów, którzy przybyli do systemu w odpowiednio długim odcinku czasu t i którzy go opuścili w tym czasie mają rozkłady normalne z wartościami oczekiwanymi i wariancjami odpowiednio λt , $\sigma_A^2 \lambda^3 t$, oraz μt , $\sigma_B^2 \mu^3 t$.

Niech N będzie zmienną losową oznaczającą liczbę zgłoszeń przebywających w systemie, oraz $p_i = P(N = i)$, gdzie $i = 0, 1, 2, \dots$. Prawdopodobieństwa te można wyznaczyć jedynie gdy system znajduje się w stanie równowagi to znaczy w przypadku gdy $\mu > \lambda$. Dla oszacowania wartości p_i użyto, aproksymacji dyfuzyjnej zastępując dyskretny proces N ciągłym procesem dyfuzji X . W procesie tym założono, że liczba zgłoszeń przebywających w systemie zależy jedynie od dwóch pierwszych momentów rozkładów odstępów czasu pomiędzy zgłoszeniami każdego z klientów do systemu oraz czasu obsługi na stanowisku. Przyjęto więc, że wyniki obliczeń będą takie same dla różnych rozkładów mających takie same dwa pierwsze momenty. Z równań procesu dyfuzji [13] otrzymano funkcję gęstości rozkładu zmiennej losowej X w postaci

$$f(x) = \begin{cases} \frac{\lambda p_o}{\mu - \lambda} (1 - e^{zx}) & 0 < x < 1 \\ \frac{\lambda p_o}{\mu - \lambda} (e^{-z} - 1) e^{zx} & x \geq 1 \end{cases} \quad (1)$$

gdzie $z = \frac{2(\lambda - \mu)}{\sigma_A^2 \lambda^3 + \sigma_B^2 \mu^3}$

Średnią liczbę zgłoszeń przebywających w systemie możemy wyznaczyć ze wzoru

$$S = \int_0^\infty x f(x) dx = \left[0,5 + \frac{\sigma_A^2 \lambda^2 \rho + \sigma_B^2 \mu^2}{2(1 - \rho)} \right] \rho \quad (2)$$

lub

$$S = \sum_{l=1}^\infty p_l l = \left[1 + \frac{\sigma_A^2 \lambda^2 \rho + \sigma_B^2 \mu^2}{2(1 - \rho)} \right] \rho \quad (3)$$

gdzie $\rho = \frac{\lambda}{\mu}$

$$p_l = \int_{l-1}^l f(x) dx, \quad l = 1, 2, \dots$$

$$p_o = 1 - \sum_{l=1}^\infty p_l$$

Wzór (3) polecany jest [23], gdy wartości $\sigma_A^2 \lambda^2$, $\sigma_B^2 \mu^2$ oraz ρ są niewielkie.

Zauważmy, że jeśli

$$\sigma_A^2 \lambda^2 = 1, \quad \sigma_B^2 \mu^2 = 1$$

to na podstawie wzoru (2) otrzymujemy $S = \frac{\rho}{1 - \rho}$ (wynik dokładny dla systemu M/M/1)

$$\sigma_A^2 \lambda^2 = 1, \quad \sigma_B^2 \mu^2 = 0$$

to na podstawie wzoru (2) otrzymujemy $S = \frac{\rho}{2(1 - \rho)}$ (wynik dokładny dla systemu M/D/1).

Przyjmijmy:

- rozważany system obsługuje K klas elementów; każda klasa k jest niezależnym źródłem zgłoszeń, w którym zgłoszenia następują w odstępach czasu o rozkładzie z wartością oczekiwaną $\frac{1}{\lambda_k}$ i wariancją σ_{kA}^2 ,
- wszystkie klasy mają ten sam rozkład odstępów czasu pomiędzy kolejnymi zgłoszeniami,
- czasy obsługi dla klientów różnych klas mają taki sam rozkład prawdopodobieństwa z wartością oczekiwaną $\frac{1}{\mu}$ i wariancją σ_B^2 .

Parametry takiego systemu możemy wyznaczyć ze wzorów

$$\lambda = \sum_{k=1}^K \lambda_k, \quad \sigma_A^2 = \sum_{k=1}^K \left(\frac{\lambda_k}{\lambda} \right)^3 \sigma_{kA}^2, \quad \rho = \frac{\sum_{k=1}^K \lambda_k}{\mu} \quad (4)$$

Średni czas przebywania zgłoszenia w systemie możemy wyznaczyć ze wzoru Little'a

$$T = \frac{S}{\lambda} \quad (5)$$

2. Konstrukcja modelu

Niech dane będą zbiory:

$I = \{1, 2, 3, \dots, n\}$ numerów obiektów w sieci,

$J = \{1, 2, 3, \dots, m\}$ numerów potencjalnych lokalizacji systemów,

$K = \{1, 2, 3, \dots, K\}$ numerów klas klientów.

Przyjmijmy oznaczenia:

t_{ij} – średni czas przejazdu pomiędzy j tą lokalizacją systemu a i -tym obiektem, $i \in I, j \in J$

$\frac{1}{\mu}$ – średni czas obsługi w systemie,

σ_B^2 – wariancja rozkładu czasu obsługi systemie,

λ_k – średnia liczba zgłoszeń k -tej klasy w jednostce czasu,

σ_{kA}^2 – wariancja rozkładu odstępów czasu pomiędzy kolejnymi zgłoszeniami z k -tej klasy,

p – maksymalna liczba zakładanych systemów (punktów serwisowych).

Za zmienne decyzyjne przyjmujemy:

$$y_j = \begin{cases} 1 & \text{gdy system zlokalizowano w } j \text{ tym obiekcie} \\ 0 & \text{w przeciwnym wypadku,} \end{cases}$$

$$x_{ij} = \begin{cases} 1 & \text{gdy } i\text{-ty obiekt obsługiwany jest przez system w} \\ 0 & \text{w przeciwnym wypadku,} \end{cases}$$

Dla zagadnienia lokalizacji centrów przyjęto następujące założenia:

1. Utworzone powinno zostać co najwyżej p -systemów

$$\sum_{j \in J} y_j \leq p \quad (6)$$

2. Każdy obiekt obsługiwany jest tylko przez jeden system obsługujący

$$\sum_{j \in J} x_{ij} = 1, \quad i \in I \quad (7)$$

3. Każdy obiekt obsługiwany jest przez założony system obsługujący

$$x_{ij} \leq y_j, \quad i \in I, j \in J \quad (8)$$

4. Założony system obsługujący jest w stanie równowagi

$$\sum_{i \in I} \sum_{k \in K} \lambda_k x_{ij} < \mu_j, \quad j \in J \quad (9)$$

Całkowity średni czas obsługi w systemie wyznaczamy ze wzoru

$$T_{ij} = (T_j + t_{ij}) x_{ij}, \quad i \in I, j \in J \quad (10)$$

gdzie: T_j (średni czas przebywania zgłoszenia w j tym systemie) wyznacza się ze wzoru (5).

W zagadnieniu lokalizacji centrów przyjęto kryterium minimalizacji sumy całkowitych średnich czasów obsługi

$$\sum_{i \in I} \sum_{j \in J} T_{ij} \rightarrow \min \quad (11)$$

4. Przykład lokalizacji systemu G/G/1 z zastosowaniem aproksymacji dyfuzyjnej

Model (6)-(11) zastosowano do znalezienia lokalizacji co najwyżej 5 punktów serwisowych z 14 proponowanych. Punkty te mają obsługiwać klientów z 20 miast. Czasy przejazdu pomiędzy potencjalnymi punktami serwisowymi a miastami klientów odczytano ze strony internetowej czas.dojazd.org/-7k. Przyjęto, że każdy punkt serwisowy jest systemem masowej obsługi typu G/G/1 z takimi samymi parametrami $\frac{1}{\mu} = 5h, \sigma_B = 2h$.

W tabeli podano wartości parametrów λ_i i σ_{iA} $i \in I$ w zależności od liczby istniejącego sprzętu w tym mieście, wyznaczone wartości λ_j, ρ_j, T_j $j \in J$ oraz rozwiązanie optymalne.

Tabela 1

Rozwiązanie optymalne zagadnienia lokalizacji systemów G/G/1

Lokalizacja Klient	Białystok	Bydgoszcz	Gdańsk	Kraków	Lublin	Łódź	Poznań	Rzeszów	Szczecin	Świecko	Terespol	Toruń	Warszawa	Wrocław	λ_i	σ_{iA}	T_{ij}
1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18
Białystok													1		0,014	0,002	8,28
Bydgoszcz		1													0,021	0,002	4,14
Cieszyn				1											0,014	0,002	7,65
Gdańsk		1													0,028	0,002	6,75
Katowice				1											0,032	0,002	6,98
Kielce				1											0,014	0,002	7,88
Końbaskowo									1						0,007	0,002	3,63
Kraków				1											0,028	0,002	5,80
Lublin													1		0,014	0,002	7,87
Łódź													1		0,028	0,002	7,33
Medyka				1											0,006	0,002	9,88
Olsztyn		1													0,013	0,002	7,21
Opole							1								0,019	0,002	8,72
Poznań							1								0,033	0,002	4,80
Rzeszów				1											0,014	0,002	8,42
Szczecin									1						0,028	0,002	3,32
Świecko									1						0,003	0,002	6,40
Toruń		1													0,008	0,002	4,83
Warszawa													1		0,042	0,002	5,25
Wrocław							1								0,035	0,002	7,42

cd. tabeli 1

1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18
μ	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2	0,2		suma	133
σ_B	2	2	2	2	2	2	2	2	2	2	2	2	2	2			
λ_j		0,1		0,1			0,1		0,04				0,1				
ρ_j		0,3		0,5			0,4		0,2				0,5				
T_j		4,1		5,8			4,8		3,3				5,3				

Przy powyższych lokalizacjach funkcja celu przyjmuje wartość 133h.

Podsumowanie

W pracy przyjęto, że każde centrum dystrybucji jest systemem masowej obsługi G/G/1. Za kryterium oceny lokalizacji przyjęto średni całkowity czas obsługi zgłoszeń (11).

Centra w tym zagadnieniu są rozmieszczone w ten sposób, aby dostosować liczbę klientów, których położenie (w sensie czasu przejazdu do centrum) jest odpowiednio bliskie centrum do możliwości obsługi tych klientów w średnio niewielkim czasie. Należy zauważyć, że w aproksymacji dyfuzyjnej wykorzystywane są jedynie dwa parametry rozkładu: wartość oczekiwana oraz wariancja, natomiast nieuwzględniony jest kształt rozkładu. Na podstawie analizy wielu rozwiązań można zauważyć, że jeśli znane są rozkłady odstępów czasu pomiędzy kolejnymi zgłoszeniami poszczególnych klientów oraz rozkłady czasów obsługi, to należy w zagadnieniu uwzględnić te rozkłady. W przypadku, jeśli zastosujemy aproksymację dyfuzyjną dla systemu G/G/1, różnice w wielkościach średnich czasów przebywania zgłoszenia w systemie zwiększają się wraz ze wzrostem wielkości $\rho = \frac{\lambda}{\mu}$, co powoduje w tych przypadkach

różnice w lokalizacjach. Wynika stąd, że aproksymację dyfuzyjną w zagadnieniach lokalizacji systemów masowej obsługi należy stosować w przypadkach, gdy nie są znane rozkłady czasów zgłoszeń i obsługi, a można wyznaczyć średni czas pomiędzy zgłoszeniami klientów, średni czas obsługi, oraz wariancje tych czasów.

Literatura

1. Batta R. (1988). Single Server Queueing-Location Models with Rejection. Transportation Science, 22, 209-216.

2. Batta R. (1989). The Stochastic Queue Median over a Finite Discrete Set. *Operations Research* 37, 648-652.
3. Batta R., Prasad S.Y. (1993). Determining Efficient Facility Locations on a Tree Network Operating as a FIFO M/G/1 Queue. *Networks*, 23, 597-603.
4. Batta R., Berman O. (1989). A Location Model for a Facility Operating as an M/G/k Queue, A Location Model for a Facility Operating as an M/G/k Queue. *Networks* 19, 717-728.
5. Batta R., Larson R.C., Odoni A.R. (1988). A Single-server Priority Queueing-location Model. A Single-server Priority Queueing-location Model. *Networks*, 18, 87-103.
6. Berman O., Larson R., Chiu S.S. (1985). Optimal Server Location on a Network Operating as an M/G/1 Queue. *Operations Research*, 33, 46-71.
7. Berman O., Larson R.C., Parkan C. (1985). The Stochastic Queue P-Median Problem. *Transportation Science*, 21, 207-216.
8. Brandeau M.L., Chiu S.S. (1990). A Unified Family of Single-server Queueing Location Models. *Operation Research*, 38, 1034-1044.
9. Brotcorne L., Laporte G., Semet F. (2003). Ambulance Location and Relocation Models. *European Journal of Operational Research*, 147, 451-463.
10. Chiu S.S., Berman O., Larson R.C. (1985). Locating a Mobile Server Queueing Facility on a Tree Network. *Management Science*, 31, 764-716.
11. Coullard C.R., Daskin M.S., Shen Z.J. (2003). A Joint Location-Inventory Model. *Transportation Science*, 37(1). (40-55).
12. Coullard C.R., Daskin M.S., Shen Z.J. (2002). An Inventory Location Model: Formulation, Solution Algorithm and Computational Results. *Annals of Operation Research*, 110, 83-106.
13. Czachórski T. (1999). Modele kolejkowe w ocenie efektywności sieci i systemów komputerowych. Pracownia Komputerowa Jacka Skalmierskiego, Gliwice.
14. Daskin M.S. (2001). A New Approach to Solving the Vertex p-center Problem to Optimality: Algorithm and Computational Results. *Communication of the Operations Research Society of Japan* 45, 428-436.
15. Jakowska-Suwalska K. (2006a). Zagadnienie lokalizacji systemów masowej obsługi z wieloma klasami klientów. *Badania operacyjne i systemowe*. W: Wiedza systemowa dla rozwoju regionów i przedsiębiorstw w Polsce. Red. J. Stachowicz, A. Straszak, S. Walukiewicz. Akademicka Oficyna Wydawnicza Exit. Warszawa, 195-202.
16. Jakowska-Suwalska K. (2006b). Zagadnienie lokalizacji systemów masowej obsługi. W: *Modelowanie preferencji a ryzyko '06*. Red. T. Trzaskalik. AE, Katowice, 225-234.

-
17. Larson R.C. (1974). A Hypercube Queuing Model for Facility Location and Redistricting in Urban Emergency Services. *Computers and Operation Research*, 1, 61-95.
 18. Mamnoon J., Baveja A., Batta R. (1999). The Stochastic Queue Center Problem. *Computer and Operation Research*, 26, 1423-1436.
 19. Marianov V., ReVelle C. (1996). The Queueing Maximal Availability Location Problem: A Model for the Siting of Emergency Vehicles. *European Journal of Operational Research* 93, 110-120.
 20. Marianov V., ReVelle C. (2004). Location-Allocation of Multiple-Server Service Centers with Constrained Queues or Waiting Times. *Annals of Operations Research*, 35-50.
 21. Marianov V., ReVelle C.S. (1996). The Queueing Maximal Availability Location Problem. A Model for Siting of Emergency Vehicles. *European Journal of Operational Research*, 93, 12-120.
 22. Marianov V., ReVelle C.S. (2003). Location Models for Airline Hubs Behaving as M/D/c Queues. *Computers and Operations Research*, 30, 983-1003.
 23. Reiser M., Kobayashi H. (1974). Accuracy of the Diffusion Approximation for Some Queueing Systems. *IBM Journal of Res. Develop.*, 18, 110-124.

Marta Kapica

Cezary Dominiak

Akademia Ekonomiczna w Katowicach

OPTIMALIZACJA SYSTEMU MASOWEJ OBSŁUGI NA PRZYKŁADZIE ŚLĄSKIEGO WESOŁEGO MIASTECZKA

Wprowadzenie

Współczesne systemy gospodarcze charakteryzują się wysokim i stale rosnącym udziałem sektora usług. Wynika to przede wszystkim z rozwoju automatyzacji procesów produkcyjnych, technologii informatycznych i zmian społecznych, które są ich rezultatem. Istotnym czynnikiem determinującym efektywność w tym sektorze gospodarki jest odpowiednia przepustowość kanałów obsługi klienta, która bezpośrednio wpływa na wielkość przychodów ze sprzedaży. Wynika to z faktu, że zasoby zaspokajające popyt w sferze usług nie mogą być magazynowane, w związku z czym powinny być one dostosowywane do przewidywanego poziomu i struktury popytu, przy czym powinny być dostosowywane w taki sposób, aby minimalizować niezbędne nakłady inwestycyjne oraz koszty operacyjne funkcjonowania systemu.

W niniejszej pracy przedstawiony został przykład optymalizacji systemu obsługi klientów w Śląskim Wesołym Miasteczku w Chorzowie. Pokazane zostało wykorzystanie teorii masowej obsługi do optymalizacji ilości stanowisk kasowych w Śląskim Wesołym Miasteczku w Chorzowie zapewniającej odpowiedni poziom obsługi klientów (czyli czas oczekiwania w kolejce nieprzekraczający maksymalnego zadanego czasu oczekiwania).

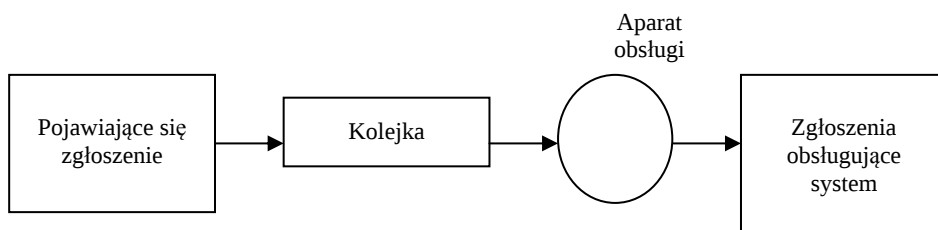
Ponieważ struktura popytu (napływu zgłoszeń) charakteryzuje się bardzo wysoką niejednorodnością (rozpatrywaną w ujęciu dziennym), to dla potrzeb optymalizacji kosztów funkcjonowania systemu (planowania ilości czynnych kanałów obsługi) wykorzystany został model programowania liniowego w liczbach całkowitych, który został zaimplementowany w arkuszu kalkulacyjnym Excel.

W pierwszej części niniejszej pracy krótko przedstawiono podstawowe informacje dotyczące teorii systemów masowej obsługi, następnie omówiono sposób funkcjonowania Śląskiego Wesołego Miasteczka, ze szczególnym uwzględnieniem systemu obsługi klientów. W trzeciej części zostały przedstawione główne założenia badawcze, omówiono sposób przeprowadzenia badania oraz pokazano wybrane fragmenty i rezultaty obliczeń. Pracę kończy podsumowanie, w którym zawarto wnioski z przeprowadzonych prac.

1. Podstawowe pojęcia teorii masowej obsługi

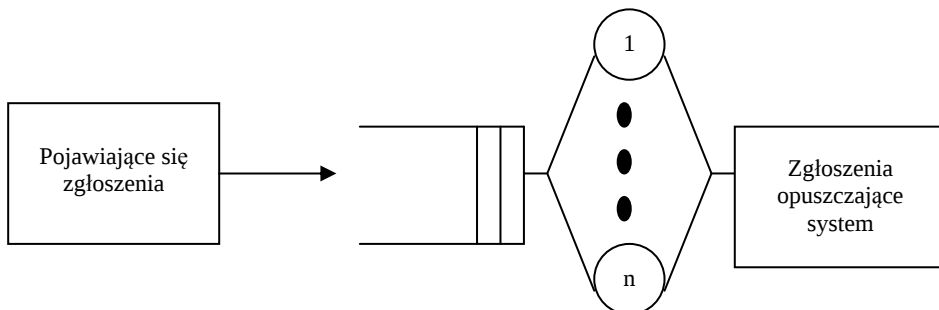
Teorię systemów masowej obsługi (zwaną również teorią kolejek) jako pierwszy sformułował Agner Krarup Erlang, obserwując prace central telefonicznych. W 1909 roku wydał pierwszą pracę *The Theory of Probabilities and Telephone Conversations*, w której udowadnia, że połączenia telefoniczne podlegają procesowi Poissona. Jednak za jego najważniejszą pracę uważa się, napisaną w 1917 roku, *Solution of some Problems in the Theory of Probabilities of Significance in Automatic Telephone Exchanges*, w której zawarł swoje klasyczne sformułowania dotyczące straty i oczekiwania. W 1953 roku powstała następna bardzo ważna praca dotycząca teorii masowej obsługi, autorstwa Davida Georga Kendalla. Zawarta w niej notacja, jaką należy stosować w celu oznaczenia systemu, używana jest do chwili obecnej. Z biegiem czasu pojawiało się wiele prac dotyczącej teorii masowej obsługi, a jej zastosowanie przeniesione zostało na wiele innych obszarów, związanych między innymi z transportem, handlem, usługami, produkcją czy też komunikacją.

Dwa proste schematy systemów masowej obsługi przedstawiono na rys. 1 i 2. Pierwszy z nich przedstawia system masowej obsługi z jednym aparatem obsługi, natomiast kolejny pokazuje schemat systemu z n aparatami obsługi.



Rys. 1. System kolejkowy z 1 aparatem obsługi

Źródło: [1].

Rys. 2. System kolejkowy z n aparatami obsługi

Źródło: Queueing Networks and Markov Chains. Gunter Bolch.

Podstawowymi charakterystykami systemów masowej obsługi są: strumień zgłoszeń, proces obsługi i regulamin kolejki. Strumień zgłoszeń jest opisem samego procesu przybywania zgłoszeń do systemu, opisywanym za pomocą funkcji rozkładu odstępów czasu pomiędzy kolejnymi zgłoszeniami. Proces obsługi to ciąg operacji, których realizacja jest niezbędna do obsługi zgłoszenia w systemie. Główną charakterystyką procesu obsługi jest czas jego trwania (pomijamy tu czas, w którym zgłoszenie oczekuje na obsługę w kolejce). Regulamin kolejki określa kolejność, w jakiej zgłoszenia znajdujące się w kolejce, są wybierane w celu obsłużenia np. FIFO, LIFO, SIRO (Selection In Random Order), RRS (Random Selection for Service), RR (Round – Robin), PS (Processor – Sharing). Ponieważ zarówno w praktyce, jak i w literaturze rozpatruje się wiele różnych typów systemów masowej obsługi, często w celu wprowadzenia systematyki stosuje się konwencję zaproponowaną przez Kendalla: $X/Y/m$, gdzie:

X – symbol rozkładu wejściowego strumienia zgłoszeń,

Y – symbol rozkładu czasów obsługi zgłoszeń,

m – liczba kanałów obsługi.

Typy rozkładów strumienia wejściowego oraz czasów obsługi oznaczamy przez następującymi symbolami:

D – strumień zdeterminowany lub regularny,

M – wykładniczy rozkład czasów obsługi lub odstępów czasu pomiędzy sąsiednimi zgłoszeniami, tzn. poissonowski rozkład przybyć,

E_k – rozkład Erlanga k -tego rzędu, który może wystąpić zarówno po stronie urządzeń obsługujących, jak i po stronie zgłoszeń,

N – rozkład normalny,

H_r – rozkład hiperwykładniczy rzędu r ,

C_k – rozkład Coxa rzędu k ,

GI – strumień ogólnego typu, dowolny i niezależny,

G – strumień o dowolnym rozkładzie czasów obsługi.

Niżej krótko scharakteryzowano dwa podstawowe modele systemów obsługi: M/M/1 oraz M/M/k. Drugi z nich został wykorzystany w dalszej części pracy.

1.1. System M/M/1 [11]

Strumień zgłoszeń jest opisem samego procesu przybywania zgłoszenia do systemu, opisywanym za pomocą funkcji rozkładu odstępów czasu pomiędzy kolejnymi zgłoszeniami, co przedstawia poniższy wzór [7]

$$\lambda = \frac{1}{t_a}$$

gdzie: λ – oznacza średnie natężenie strumienia zgłoszeń, t_a – średnią długość odstępów czasu pomiędzy dwoma następującymi po sobie zgłoszeniami.

System ten składa się z jednego aparatu obsługi, a odstęp czasu pomiędzy dwoma kolejnymi zgłoszeniami jest zmienną o rozkładzie wykładniczym

$$f(t) = \lambda e^{-\lambda t} \quad \text{dla } t \geq 0$$

Liczba zgłoszeń (n) w systemie, w jednostce czasu o długości T jest zmienną losową o rozkładzie Poissona

$$P\{n = k\} = \frac{(\lambda T)^k e^{-\lambda T}}{k!} \quad \text{dla } k=0, 1, 2, \dots$$

Czas obsługi zgłoszenia jest zmienną losową o rozkładzie wykładniczym

$$g(t) = \mu e^{-\mu t} \quad \text{dla } t \geq 0$$

Parametr intensywności obliczany jest na podstawie wzoru

$$\rho = \frac{\lambda}{\mu}$$

Prawdopodobieństwo, że w systemie nie ma zgłoszeń obliczany jest w następujący sposób

$$P_0 = 1 - \frac{\lambda}{\mu}$$

Średnią liczbę zgłoszeń w kolejce wyznacza wzór

$$L_q = \frac{\lambda^2}{\mu(\mu - \lambda)}$$

Do obliczenia średniej liczby zgłoszeń w systemie wykorzystuje się wzór

$$L = L_q \frac{\lambda}{\mu}$$

Aby wyznaczyć średni czas spędzany przez zgłoszenie w kolejce, posługujemy się zależnością

$$W_q = \frac{L_q}{\lambda}$$

Średni czas spędzany przez zgłoszenie w systemie wyznacza wzór

$$W = W_q + \frac{1}{\mu}$$

Wzór określający prawdopodobieństwo, że zgłoszeń czekających na obsługę jest więcej niż n_0 jest następujący

$$P_{n>n_0} = \left(\frac{\lambda}{\mu}\right)^{n_0+1}$$

Prawdopodobieństwo, że w systemie będzie n zgłoszeń wyznaczane jest w następujący sposób

$$P_n = \left(\frac{\lambda}{\mu}\right)^n P_0$$

Do wyznaczenia oczekiwanego czasu przestoju kanału obsługi (w okresie czasu $[0, T]$) służy wzór

$$W_T = T \left(1 - \frac{\lambda}{\mu}\right)$$

Oczekiwany czas zajętości kanału obsługi (w okresie czasu $[0, T]$) obliczane jest na podstawie wzoru

$$B_T = T \frac{\lambda}{\mu}$$

W opisanym wyżej systemie M/M/1, analizę systemu rozpoczyna się od zbadania zależności pomiędzy stopą zgłoszeń a stopą obsługi, czyli weryfikacji tego, czy system jest stabilny. Jeżeli stopa zgłoszeń będzie mniejsza od stopy obsługi, czyli $\lambda < \mu$, oznacza to, iż układ jest stabilny. Natomiast w sytuacji, gdy stopa zgłoszeń byłaby większa od stopy obsługi, należy rozważyć możliwość uruchomienia kolejnego aparatu obsługi, ponieważ układ jest niestabilny, co oznacza, iż w każdej następnej jednostce czasu długość kolejki będzie rosła.

1.2. System M/M/k [11]

Rozpatrywany system M/M/k różni się od systemu M/M/1 tylko liczbą kanałów, które działają niezależnie od siebie i są jednakowe. W systemie wielokanałowym, intensywność ruchu wyznacza wzór

$$\rho = \frac{\lambda}{k\mu}$$

Prawdopodobieństwo, że w systemie nie ma zgłoszeń obliczane jest następująco

$$P_0 = \frac{1}{\sum_{n=0}^{k-1} \frac{(\lambda/\mu)^n}{n!} + \frac{(\lambda/\mu)^k}{k!} \left(\frac{k\mu}{k\mu - \lambda} \right)}$$

Do wyznaczenia średniej liczby zgłoszeń w kolejce korzysta się z zależności

$$L_q = \frac{(\lambda/\mu)^k \lambda\mu}{(k-1)!(k\mu - \lambda)^2} P_0$$

Wzór pozwalający na wyznaczenie średniej liczby zgłoszeń w całym systemie to

$$L = L_q + \frac{\lambda}{\mu}$$

Średni czas spędzony przez zgłoszenie w kolejce obliczany jest na podstawie wzoru

$$W_q = \frac{L_q}{\lambda}$$

Z kolei do oszacowania średniego czasu spędzonego przez zgłoszenie w systemie służy wzór

$$W = W_q + \frac{1}{\mu}$$

Prawdopodobieństwo, że zgłoszenie musi czekać na obsługę obliczane jest na podstawie wzoru

$$P_w = \frac{1}{k!} \left(\frac{\lambda}{\mu} \right)^k \left(\frac{k\mu}{k\mu - \lambda} \right) P_0$$

W systemie tym zakłada się powstanie jednej kolejki, która jest wspólna dla wszystkich aparatów obsługi, a powstaje w momencie, gdy wszystkie aparaty obsługi są zajęte. Podobnie jak w systemach jednokanałowych, tak również tutaj bardzo istotne zbadanie jest zależności pomiędzy liczbą zgłoszeń napływających do systemu a liczbą zgłoszeń obsługiwanych w określonej jednostce czasu. Jeżeli prawdziwa jest nierówność $\lambda < \mu$, zdolności kanału są wystarczające, oznacza to, że jeden kanał obsługi jest zdolny do obdłużenia większej ilości zgłoszeń niż ta, która pojawiła się w systemie. W systemach

wielokanałowych zależność ta liczona jest głównie przy uwzględnieniu liczby czynnych aparatów obsługi czyli gdy $\lambda < k * \mu$, zdolność obsługi systemu jest wystarczająca. Sytuacja pogarsza się, jeżeli do systemu w określonej jednostce czasu przybywa się więcej zgłoszeń niż kanały są w stanie obsłużyć (niestabilność systemu). W takiej sytuacji należy rozważyć, czy istnieje możliwość zwiększenia intensywności obsługi lub zwiększenia liczby kanałów obsługi.

Najogólniej rzecz ujmując celem analizy systemów masowej obsługi obejmuje poszukiwanie odpowiedzi na następujące problemy:

1. Czy system jest stabilny?
2. Jeśli system nie jest stabilny, to o ile należy zwiększyć liczbę kanałów obsługi lub o ile należy skrócić czas obsługi, aby uzyskać stabilność systemu?
3. Jakie jest prawdopodobieństwo, że w systemie nie ma żadnych zgłoszeń?
4. Jaka jest średnia liczba zgłoszeń w kolejce?
5. Jaka jest średnie liczba zgłoszeń w systemie (czyli oczekujących w kolejce i obsługiwanych przez kanał obsługi)
6. Jaki jest średni czas oczekiwania w kolejce?
7. Jaki jest średni czas pobytu zgłoszenia w systemie obsługi (czyli łącznie czas spędzony w kolejce oraz w obsłudze)?
8. Jakie jest prawdopodobieństwo, że nadchodzące zgłoszenie będzie musiało czekać na obsługę?
9. Jakie jest prawdopodobieństwo, że w systemie będzie n zgłoszeń?

2. Śląskie Wesołe Miasteczko w Chorzowie

Śląskie Wesołe Miasteczko w Chorzowie znajduje się na terenie Wojewódzkiego Parku Kultury i Wypoczynku im. Gen. Jerzego Ziętka (WPKiW). Chorzowski park to jedyny w Polsce tego rodzaju obiekt przyrodniczo-rekreacyjny, stworzony dla celów wypoczynku masowego. Park powstał w centrum Górnośląskiego Okręgu Przemysłowego, w miejscu, gdzie wcześniej znaleźć można było jedynie liczne hałdy i nieużytki. Rozwijający się w błyskawicznym tempie przemysł, zanieczyszczone powietrze oraz wody, brak zieleni w pobliżu miast sprawiły, że mieszkańcy śląska zapragnęli mieć swój własny park, duży, zielony, bogato wyposażony w urządzenia mające służyć wypoczynkowi i rozrywce. Powołana w 1951 roku Komisja Planu Regionalnego Górnośląskiego Okręgu Przemysłowego (GOP), w celu przekształcenia struktury gospodarczej i przestrzennej oraz racjonalnego zagospodarowania przestrzennego Górnego Śląska i Zagłębia, za jeden z głównych celów postawiła sobie wybudowanie Wojewódzkiego Parku Kultury i Wypoczynku. Już 20 grudnia 1950 roku, na posiedzeniu Wojewódzkiej Rady Narodowej w Katowicach, zapadła odpowied-

nia uchwała, w maju 1951 roku Prezydium Rządu zatwierdziło budowę Parku, a w lipcu powołano do życia Komitet Budowy [12]. Na jego czele stanął generał Jerzy Ziętek, inicjator idei powstania „zielonych płuc” dla Śląska. Profesor Władysław Niemirski, projektant parku, swój projekt powiązał z uwarunkowaniami topologicznymi obszaru, na którym park miał powstać, dzięki czemu chorzowski park to nie tylko liczne ścieżki i alejki, ale również mniejsze i większe stawy, oczka wodne oraz wielkie zlewisko wodne. Prace nad przywróceniem życia środowisku tak bardzo zniszczonemu przemysłem od początku cieszyły się powszechnym zainteresowaniem i dużym poparciem mieszkańców, którzy chętnie oferowali swoją pomoc przy sadzeniu drzew, krzewów oraz wykonywaniu robót ziemnych [14]. Z biegiem lat na terenie parku powstawały nowe obiekty oraz zagospodarowane tereny zielone. Do osiągnięcia tego celu przemieszczono ok. 3,5 mln m³ ziemi, dowieziono 0,5 mln m³ torfu, a w pierwszych latach budowy wykorzystano ponad 3,5 mln drzew i krzewów aby stworzyć miejsce, o którym mieszkańcy śląska od tak dawna marzyli [15]. Obecnie zajmujący ponad 600 hektarów park, będący jednym z największych tego typu obiektów w Europie, podzielić można na dwa typy. Pierwszy z nich to typ parkowo-leśny, przeznaczony dla osób pragnących spędzić czas w zaciszu zieleni i kwiatów. Różnego rodzaju drzewa, krzewy, trawniki oraz kwiaty pokrywają ponad 400 hektarów parku. Dla miłośników zieleni organizowane są wystawy kwiatów oraz ciesząca się niezwykłą popularnością Parada Kwiatów. Urok parku podkreśla również znajdujące się na powierzchni 5 ha rosarium, uważane za jedno z 35 najpiękniejszych ogrodów różanych na świecie. To tutaj każdego roku zachwyca odwiedzających park ponad 30 tys. róż, dodatkowo upiększonych kwitnącymi w oczkach wodnych liliami. Drugie oblicze parku ma charakter rozrywkowy i dydaktyczny, tak więc przeznaczony dla osób preferujących aktywną formę wypoczynku. To dla nich przygotowano wielokilometrowe ścieżki spacerowe, ponad 30 km ścieżek rowerowych, wypożyczalnie sprzętów wodnych oraz liczne atrakcje, mające na celu nie tylko bawić, ale również uczyć. Coroczna wystawa postępu technicznego i racjonalizacji pracy, liczne wystawy fotografii, malarstwa i rzeźby, wyrobów przemysłu ludowego, książek cieszą się dużym zainteresowaniem zarówno młodszych jak i starszych gości parku, a to wszystko ma miejsce w Hali wystawowej Kapelusz. Wśród atrakcji, które zwiedzającym oferuje park należy wymienić:

- Śląskie Wesołe Miasteczko,
- Górnośląski Park Etnograficzny,
- kąpielisko FALA,
- korty tenisowe oraz sztuczna plaża,
- park linowy Palenisko,
- kolejkę szynową,

- leżący w samym sercu parku Śląski Ogród Zoologiczny,
- najstarsze i największe w Polsce Planetarium o 23-metrowej kopule stanowiącą ekran sztucznego nieba, umiejscowione na najwyższym punkcie parku,
- Stadion Śląski,
- paintball,
- jaskinię solno-jodową.

Śląskie Wesołe Miasteczko, powstałe w 1959 roku, to jedyny tego typu stały obiekt w Polsce, zajmujący ponad 20 hektarów terenu parkowego. Zlokalizowane jest w południowo-wschodniej części parku. Usytuowane na obrzeżach malowniczego stawu, wokół którego rozmieszczonych zostało ponad 50 atrakcji, stanowi jeden z jego najpiękniejszych zakątków parku. To wszystko czyni obiekt nie tylko największym, ale również najbardziej atrakcyjnym lunaparkiem w Polsce. Wśród atrakcji śląskiego Wesołego Miasteczka każdy znajdzie coś dla siebie.

ŚWM od zawsze było najważniejszym źródłem dochodu Wojewódzkiego Parku Kultury i Wypoczynku. W zeszłym sezonie liczna odwiedzających to miejsce sięgnęła prawie 240 tys. osób. Jednakże w ubiegłych latach frekwencja w lunaparku zmalała, co było w dużej mierze spowodowane wysokim stopniem zużycia posiadanych urządzeń rozrywkowych. Aby zwiększyć atrakcyjność Śląskiego Wesołego Miasteczka, zaczęto sprowadzać nowe urządzenia. Największym wydarzeniem było niewątpliwie sprowadzenie w zeszłym roku z Anglii Rollercoastera, którego wysokość sięga 21 m, co czyni go największą kolejką w Polsce i który był niewątpliwie największą atrakcją lunaparku. Sprowadzenie kolejki zahamować miało spadek liczby klientów Wesołego Miasteczka i tak też się stało. Trafność inwestycji zachęciła do dalszego rozwoju lunaparku. Efektem tych działań jest ciągle uruchamianie nowych karuzel, głównie takich, które gwarantują użytkownikom ekstremalne doznania. Obecnie ludzie coraz częściej stawiają na aktywną formę wypoczynku, dlatego też Wesołe Miasteczko, chcąc umocnić swoją pozycję na rynku, wychodzi na przeciw ich wymaganiom oraz panującym trendom.

Do 2006 roku w Śląskim Wesołym Miasteczku obowiązywał system kar-netowy. W 2006 roku został ogłoszony przetarg na opracowanie, zainstalowanie i uruchomienie elektronicznego systemu sprzedaży i rozliczania usług Śląskiego Wesołego Miasteczka w Chorzowie. System ten miał umożliwić wykorzystanie elektrycznych kart i biletów do pobierania opłat za korzystanie z atrakcji Wesołego Miasteczka. Wprowadzenie nowego systemu kasowego wiązało się przede wszystkim z dostarczeniem niezbędnego sprzętu obejmującego między innymi:

1) jednostkę centralną umożliwiającą między innymi:

- ładowanie i doładowywanie kart i biletów,
- pobieranie opłat za korzystanie z poszczególnych atrakcji,
- archiwizację zdarzeń, danych oraz tworzenie raportów,

- 2) wyposażenie stanowisk kasowych pozwalające na:
 - doładowanie kart i biletów dowolną ilością punktów,
 - anulowanie nieprawidłowych transakcji,
 - wykonywanie raportów oraz wiele innych,
- 3) drukarek do wystawiania faktur i raportów,
- 4) mobilnych terminali,
- 5) czytników kart i biletów,
- 6) kart elektronicznych.

Nowy system rozliczeniowy Śląskiego Wesołego Miasteczka, miał swój debiut w sezonie 2007. W kasach przed wejściem należy zakupić karnet elektroniczny o minimalnej wartości 25 zł, co daje klientom równocześnie 25 punktów, ponieważ cena 1 punktu to 1 zł. Za wejście na teren Wesołego Miasteczka oraz za korzystanie z urządzeń rozrywkowych pobierana jest z karty określona ilość punktów, przedstawiona szczegółowo w załączniku. Na terenie Wesołego Miasteczka istnieje 6 punktów, umożliwiających doładowanie karty.

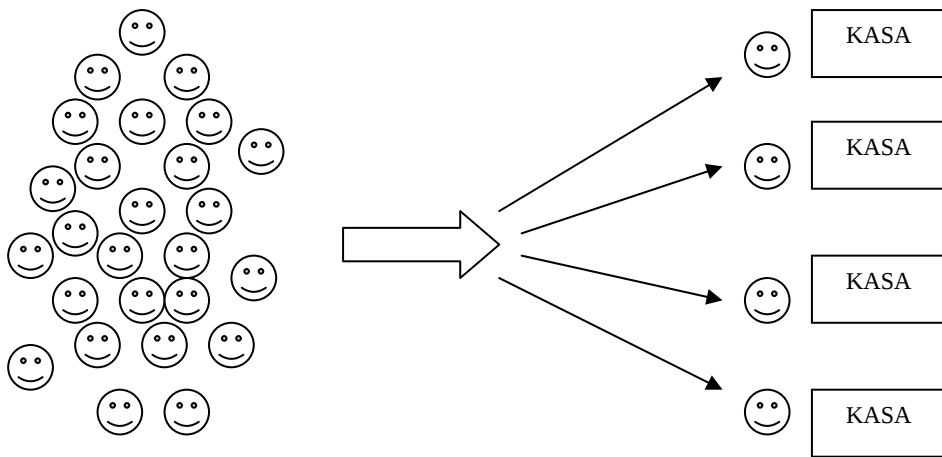
3. Analiza systemu masowej obsługi w Śląskim Wesołym Miasteczku

W tej części opracowania przedstawione zostały wyniki prac, których celem była optymalizacja systemu obsługi klientów w Śląskim Wesołym Miasteczku. Omówiono kolejno główne cechy systemu oraz model, który przyjęty został do dalszej części analizy, następnie przedstawione zostały główne aspekty dotyczące funkcjonującego systemu oraz wynikające z nich cele dalszych badań. W dalszej kolejności przedstawiono rezultaty analizy stabilności systemu, oszacowano minimalną ilość kas niezbędną do zapewnienia odpowiedniego poziomu obsługi i w ostatniej części przedstawiono model optymalizacji obsady stanowisk kasowych.

3.1. System masowej obsługi w Śląskim Wesołym Miasteczku

Niżej przedstawiono podstawowe zasady funkcjonowania systemu sprzedaży biletów w Śląskim Wesołym Miasteczku. Jako zgłoszenie traktujemy pojawienie się każdej osoby pragnącej kupić bilet, uprawniający do wejścia na jego teren. Obsługa polega na sprzedaży elektronicznego karnetu w jednej z pięciu kas. Urządzeniami obsługującymi przybywające zgłoszenia, czyli ludźmi, są kasy biletowe. System jest systemem z oczekiwaniem, ponieważ w chwili gdy wszystkie kasy są zajęte tworzy się kolejka. Mimo iż kolejka ma nieograniczoną liczbę miejsc, system ten jest systemem ze stratami. Potencjalni klienci, gdy widzą, że kolejka jest duża często rezygnują z zakupu biletu. Proces obsługi jest jednofazowy, ponieważ cała obsługa polega jedynie na sprzedaży biletu, a cały system jest systemem wielokanałowym ponieważ lunapark posiada 5 kas

biletowych. Aparaty obsługi śląskiego lunaparku są równouprawnione, co oznacza, że klient może podejść do którejkolwiek wolnej kasy biletowej w celu zakupu elektronicznego karnetu, a organizacja systemu jest nieuporządkowana. Jeżeli chodzi o regulamin kolejki, obsługa następuje według kolejności FIFO, czyli pierwsza osoba w kolejce zostaje jako pierwsza kierowana do obsługi w chwili zwolnienia aparatu. Schemat systemu sprzedaży biletów przedstawiono na rys. 3.



Rys.3. Napływ klientów do Wesołego Miasteczka

Tak więc w dalszej części analizy wykorzystany zostanie model M/M/k, przy czym w pierwszej kolejności zweryfikowane zostaną założenia odnośnie rozkładów strumienia zgłoszeń oraz obsługi.

W tym miejscu należy zwrócić uwagę na kilka istotnych czynników funkcjonowania systemu, które wynikały ze wstępnych analiz i nieformalnej obserwacji jego funkcjonowania. Pierwszym z nich jest wysoka sezonowość liczby odwiedzających, gdzie szczyt ilości osób odwiedzających przypada na miesiąc sierpień. Ponadto, istnieje wysoka sezonowość w poszczególnych dniach tygodni, gdzie szczyty przypadają na niedziele i dni świąteczne. Kolejnym niezmiernie istotnym czynnikiem jest wysoka nieregularność ilości odwiedzających w kolejnych godzinach funkcjonowania obiektu (nieregularny napływ zgłoszeń w poszczególnych porach dnia). Zaobserwowano również tworzenie się długich kolejek w godzinach porannych i południowych, przede wszystkim w niedziele i w dni świąteczne. Ustalono także, że nie ma możliwości skrócenia czasu obsługi ze względu na ograniczenia techniczne procesu sprzedaży (ładowanie karty, wydruk paragonu). Ponadto należy podkreślić, że otwarcie nowych punktów kasowych pociągałoby za sobą wysokie nakłady, gdyż wymagałoby to budowy nowych budynków z kasami oraz odpowiedniego ich zabezpieczenia i wyposażenia.

3.2. Cel analizy i procedura badawcza

Biorąc pod uwagę przedstawione w poprzednim punkcie aspekty sformułowane zostały sformułowane następujące cele analizy:

- ocenić aktualny stan systemu, w tym średni czas oczekiwania na obsługę,
- przyjmując hipotezę o niestabilności systemu ustalić minimalną liczbę stanowisk kasowych zapewniających stabilność w dniach o najwyższym natężeniu zgłoszeń,
- zbudować model optymalizacyjny umożliwiający minimalizację kosztów funkcjonowania systemu przy zapewnieniu odpowiedniego poziomu obsługi (maksymalnego czasu oczekiwania).

Aby zrealizować tak sformułowane cele badania, opracowana została następująca procedura badawcza. W pierwszym etapie przeprowadzono analizę ilości klientów na podstawie danych historycznych, aby ustalić dni, w których można spodziewać się największej liczby odwiedzających i następnie, aby w tych dniach zebrać szczegółowe dane dotyczące funkcjonowania systemu za pomocą obserwacji i pomiaru w terenie. Następnie zweryfikowano założenia o typach rozkładów, w celu weryfikacji możliwości wykorzystania modelu M/M/k. W etapie trzecim dokonano oceny stabilności systemu oraz przeprowadzono symulację zwiększenia ilości stanowisk kasowych. Następnie opracowano model optymalizacyjny do planowania harmonogramu uruchamiania stanowisk kasowych.

3.3. Rezultaty optymalizacji systemu masowej obsługi

Analiza sprzedaży i wytypowanie dni do badania

W tabeli 1 przedstawione zostały dane historyczne dotyczące ilości osób odwiedzających w podziale na kolejne miesiące sezonu oraz dni tygodnia.

Tabela 1

Średnia liczby gości Śląskiego Wesołego Miasteczka w 2006 roku

Dzień tygodnia	Średnia liczba gości						
	kwiecień	maj	czerwiec	lipiec	sierpień	wrzesień	październik
Poniedziałek	534	218	456	848	1087	154	99
Wtorek	0	321	830	1019	1674	98	184
Środa	79	407	669	629	1454	85	100
Czwartek	10	311	1465	1253	1757	41	151
Piątek	81	624	1065	995	1264	112	47
Sobota	723	788	2250	2691	3358	1071	207
Niedziela	1266	2266	3165	5669	7001	2777	315

Widać wyraźnie, że największa średnia liczba osób odwiedzających zaobserwowana została w miesiącu sierpniu w niedzielę. Tak więc podjęto decyzję, aby badanie przeprowadzić w niedzielę i dni świąteczne w sierpniu 2007, czyli w dniach 05.08.2007, 12.08.2007, 15.08.2007, 19.08.2007, 26.08.2007.

W tabeli 2 przedstawiono zebrane dane opisujące liczbę zgłoszeń w kolejnych godzinach funkcjonowania systemu w dniach wytypowanych do badania.

Tabela 2

Dane dotyczące liczby zgłoszeń w analizowanych dniach

Dzień tygodnia	godzina data	Napływ w kolejnych godzinach							
		10 ⁰⁰	11 ⁰⁰	12 ⁰⁰	13 ⁰⁰	14 ⁰⁰	15 ⁰⁰	16 ⁰⁰	17 ⁰⁰
Niedziela	05-08-2007	1244	1061	1178	1110	1013	793	793	161
Niedziela	12-08-2007	591	804	674	692	627	640	264	2
Środa (święto)	15-08-2007	1328	1172	1283	1151	1105	893	901	316
Niedziela	19-08-2007	1329	1179	1145	1115	1069	721	585	198
Niedziela	26-08-2007	1367	1216	1098	1109	1107	612	576	243

Weryfikacja założeń modelu

Następnie na podstawie zebranych danych zweryfikowano założenia modelu M/M/k. Jako hipotezę zerową przyjęto, że rozkład zgłoszeń jest rozkładem Poissona, hipotezę alternatywną, że rozkład zgłoszeń nie jest rozkładem Poissona. Odpowiednie dane przedstawiono w tabeli 3.

Tabela 3

Rezultaty testu Kołmogorowa-Smirnowa

		Sierpień 05	Sierpień 12	Sierpień 15	Sierpień 19	Sierpień 26
Parametr rozkładu Poissona(a,b)	Średnia	1027,4286	613,1429	1119,0000	1064,1429	1002,0000
Największe różnice	Wartość bezwzględna	0,423	0,428	0,398	0,509	0,491
	Dodatnia	0,286	0,143	0,286	0,286	0,286
	Ujemna	-0,423	-0,428	-0,398	-0,509	-0,491
Z Kołmogorowa-Smirnowa		1,119	1,132	1,054	1,347	1,299
Istotność asymptotyczna (dwustronna)		0,163	0,154	0,216	0,053	0,069

Na przyjętym poziomie istotności nie ma podstaw do odrzucenia hipotezy, że rozkład zgłoszeń jest rozkładem Poissona.

W rezultacie obserwacji pracy punktów kasowych ustalono, że średni czas obsługi jednego klienta wynosi ok. 18 sekund, tak więc każde stanowisko może obsłużyć średnio 200 osób w ciągu godziny.

Ocena efektywności systemu

Niżej przedstawiono dane zebrane dla 05.08 opisujące napływ klientów oraz średni czas oczekiwania w kolejce w poszczególnych porach dnia.

Tabela 4

Średni czas oczekiwania w kolejce w poszczególnych godzinach

Godzina	Średni czas oczekiwania w kolejce (w minutach)	Napływ
10-11	8	1244
11-12	11	1061
12-13	20	1178
13-14	20	1110
14-15	19	1013
15-16	10	793
16-17	4	793
17-18	1	161
18-19	0	98
19-20	0	28
20-21	0	7
21-22	0	0
22-23	0	0

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1	Godzina	10	11	12	13	14	15	16	17	18	19	20	21	22
2	λ	1244	1061	1178	1110	1013	793	793	161	98	28	7	0	0
3	μ	200	200	200	200	200	200	200	200	200	200	200	200	200
4	k	5	5	5	5	5	5	5	5	5	5	5	5	5
5														
6	Układ stabilny	nie	nie	nie	nie	nie	tak	tak	tak	tak	tak	tak	tak	tak
7														
8	p	1,244	1,061	1,178	1,110	1,013	0,793	0,793	0,161	0,098	0,028	0,007	0,000	0,000
9														
10														
11														
12														
13														
14														
15														
16														
17														
18														
19														
20														
21														
22														
23														
24														
25														
26														
27														
28														

Rys. 4. Badanie stabilności systemu

Na podstawie tabeli oraz rezultatów badania stabilności systemu pokazanych na rys. 4 widać, że do godziny 15 system jest systemem niestabilnym (zwiększa się stale czas oczekiwania na obsługę). Podobne rezultaty otrzymano dla pozostałych dni, w których przeprowadzono badanie.

Symulacja zwiększenia ilości punktów kasowych

Następnie przeprowadzono symulację stabilności systemu zwiększając ilość stanowisk kasowych do 6, a następnie do 7. Przy siedmiu stanowiskach kasowych udało się uzyskać stabilność systemu. Na rys. 5 i 6 oraz w tabeli 5 przedstawiono obliczenia przeprowadzone dla 05.08.2007. Podobne rezultaty otrzymano dla pozostałych dni, w których przeprowadzono badanie: system osiągał stabilność przy siedmiu kasach.

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1 Godzina		10	11	12	13	14	15	16	17	18	19	20	21	22
2 λ		1244	1061	1178	1110	1013	793	793	161	98	28	7	0	0
3 μ		200	200	200	200	200	200	200	200	200	200	200	200	200
4 k		7	7	7	7	7	7	7	7	7	7	7	7	7
5														
6 Układ stabilny		tak	tak	tak	tak	tak	tak	tak	tak	tak	tak	tak	tak	tak
7														
8 p		0,889	0,758	0,841	0,793	0,724	0,566	0,566	0,115	0,070	0,020	0,005	0,000	0,000
9														
10														
11														
12														
13														
14														
15														
16														
17														
18														
19														
20														
21														
22														
23														
24														
25														
26														
27														
28														

Rys. 5. Badanie stabilności systemu dla 7 stanowisk kasowych

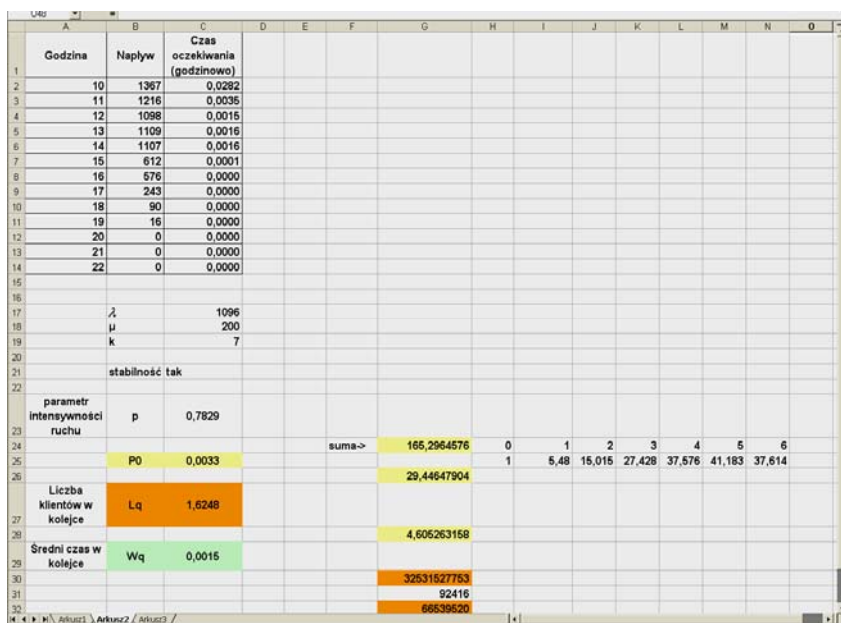
Tabela 5

Średni czas oczekiwania w kolejce po wprowadzeniu 2 nowych kas

Godzina	Napływ	Średni czas oczekiwania (w godz.)
1	2	3
10-11	1244	0,0044
11-12	1061	0,0012
12-13	1178	0,0026

cd. tabeli 5

1	2	3
13-14	1110	0,0016
14-15	1013	0,0009
15-16	793	0,0020
16-17	793	0,0020
17-18	161	0,0000
18-19	98	0,0000
19-20	28	0,0000
20-21	7	0,0000
21-22	0	0,0000
22-23	0	0,0000



Godzina	Napływ	Czas oczekiwania (godzinowo)															
10	1367	0,0282															
11	1216	0,0035															
12	1098	0,0015															
13	1109	0,0016															
14	1107	0,0016															
15	612	0,0001															
16	576	0,0000															
17	243	0,0000															
18	90	0,0000															
19	16	0,0000															
20	0	0,0000															
21	0	0,0000															
22	0	0,0000															
λ		1096															
μ		200															
k		7															
stabilność tak																	
parametr intensywności ruchu	p	0,7829															
suma->			165,2964576	0	1	2	3	4	5	6							
P0	0,0033			1	5,48	15,015	27,428	37,576	41,183	37,614							
Liczba klientów w kolejce	Lq	1,6248															
Średni czas w kolejce	Wq	0,0015															
			32531527753														
			52416														
			66539520														

Rys. 6. Arkusz wykorzystany do wykonania obliczeń

Model optymalizacyjny do harmonogramowania pracy kas

Niżej przedstawiono model optymalizacyjny do harmonogramowania pracy kas. Ponieważ wykorzystanie modelu ma sens jedynie w te dni, gdy system jest stabilny, przedstawiony został dla danych zebranych 14.07.2007 (sobota).

Badanie rozpoczęto od przedstawienia, jak kształtował się napływ klientów w analizowanym dniu oraz określenia, jaka jest wymagana liczba czynnych kas w każdej godzinie, zapewniająca optymalny serwis.

Tabela 6

Dane dotyczące napływu klientów 14 lipca

Napływ klientów dla dnia 14.07.2007											
10-11	11-12	12-13	13-14	14-15	15-16	16-17	17-18	18-19	19-20	20-21	21-22
545	642	735	554	431	275	257	154	62	27	0	0

Możliwość obsługi każdej kasy wynosi 200 osób w ciągu godziny, tak więc wymagana ilość otwartych kas przedstawia się następująco:

Tabela 7

Dane dotyczące wymaganej ilości otwartych kas

Wymagana liczba otwartych punktów obsługi											
10-11	11-12	12-13	13-14	14-15	15-16	16-17	17-18	18-19	19-20	20-21	21-22
3	4	4	3	3	2	2	1	1	1	1	1

Zmienne decyzyjne:

X_1 – ilość kas otwieranych o godzinie 10⁰⁰

X_2 – ilość kas otwieranych o godzinie 11⁰⁰

X_3 – ilość kas otwieranych o godzinie 12⁰⁰

X_4 – ilość kas otwieranych o godzinie 13⁰⁰

X_5 – ilość kas otwieranych o godzinie 14⁰⁰

X_6 – ilość kas otwieranych o godzinie 15⁰⁰

X_7 – ilość kas otwieranych o godzinie 16⁰⁰

X_8 – ilość kas otwieranych o godzinie 17⁰⁰

X_9 – ilość kas otwieranych o godzinie 18⁰⁰

X_{10} – ilość kas otwieranych o godzinie 19⁰⁰

X_{11} – ilość kas otwieranych o godzinie 20⁰⁰

X_{12} – ilość kas otwieranych o godzinie 21⁰⁰

Funkcja celu

Zadanie polega na minimalizacji funkcji celu, przedstawiającej ilość kas potrzebnych do obsłużenia przybyłych gości.

$$F(X) \sum_{n=1}^{12} X_n \rightarrow \min$$

Warunki ograniczające zadania przedstawiają się następująco

$$\begin{aligned}
 X_1 &\geq 3 \\
 X_1 + X_2 &\geq 4 \\
 X_1 + X_2 + X_3 &\geq 4 \\
 X_1 + X_2 + X_3 + X_4 &\geq 3 \\
 X_1 + X_2 + X_3 + X_4 + X_5 &\geq 3 \\
 X_1 + X_2 + X_3 + X_4 + X_5 + X_6 &\geq 2 \\
 X_1 + X_2 + X_3 + X_4 + X_5 + X_6 + X_7 &\geq 2 \\
 X_1 + X_2 + X_3 + X_4 + X_5 + X_6 + X_7 + X_8 &\geq 1 \\
 X_2 + X_3 + X_4 + X_5 + X_6 + X_7 + X_8 + X_9 &\geq 1 \\
 X_3 + X_4 + X_5 + X_6 + X_7 + X_8 + X_9 + X_{10} &\geq 1 \\
 X_4 + X_5 + X_6 + X_7 + X_8 + X_9 + X_{10} + X_{11} &\geq 1 \\
 X_5 + X_6 + X_7 + X_8 + X_9 + X_{10} + X_{11} + X_{12} &\geq 1
 \end{aligned}$$

$$X_i \geq 0, i = 1, 2, \dots, 12$$

$$X_i \in C, i = 1, 2, \dots, 12$$

Na rys. 7 i 8 pokazano zapis modelu w arkuszu kalkulacyjnych Excel oraz otrzymane rezultaty.

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q
	Pora dnia	Liczba otwieranych kas	Konieczna liczba kas	Liczba czynnych kas	Napływ klientów dla dnia 14-07-2007												
1	10-11	0	3	0	545												
2	11-12	0	4	0	642												
3	12-13	0	4	0	735												
4	13-14	0	3	0	554												
5	14-15	0	3	0	431												
6	15-16	0	2	0	275												
7	16-17	0	2	0	257												
8	17-18	0	1	0	154												
9	18-19	0	1	0	62												
10	19-20	0	1	0	27												
11	20-21	0	1	0	0												
12	21-22	0	1	0	0												
13	suma	0															

Rys. 7. Arkusz wykorzystany do optymalnego harmonogramowania pracy kas

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q
1		Liczba otwieranych kas	Konieczna liczba kas	Liczba czynnych kas	Napływ klientów dla dnia 14-07-2007												
2	10 – 11	4	3	4	545												
3	11 – 12	0	4	4	642												
4	12 – 13	0	4	4	735												
5	13 – 14	0	3	4	554												
6	14 – 15	1	3	5	431												
7	15 – 16	0	2	5	275												
8	16 – 17	0	2	5	267												
9	17 – 18	0	1	5	154												
10	18 – 19	0	1	1	62												
11	19 – 20	0	1	1	27												
12	20 – 21	0	1	1	0												
13	21 – 22	0	1	1	0												
14	suma	5															

Rys. 8. Arkusz z rozwiązaniem optymalnym

Rozwiązanie optymalne prezentuje się następująco: $X_1 = 4$, $X_2 = 0$, $X_3 = 0$, $X_4 = 0$, $X_5 = 1$

$X_6 = 0$, $X_7 = 0$, $X_8 = 0$, $X_9 = 0$, $X_{10} = 0$, $X_{11} = 0$, $X_{12} = 0$. Wartość funkcji celu wyniosła 5.

Przedstawiony przykład wykorzystywał dane historyczne, praktyczne jego wykorzystanie wymagałoby ustalania prognozowanej ilości osób odwiedzających oraz struktury napływu na kolejny dzień i na tej podstawie ustalanie harmonogramu otwierania kas.

Podsumowanie

Przedstawione w niniejszym opracowaniu rezultaty badań pokazują, że wykorzystanie teorii masowej obsługi w połączeniu z optymalizacją liniową mogą być skutecznym narzędziem podnoszącym efektywność funkcjonowania systemów masowych.

Zaprezentowany przykład pokazuje, że w Śląskim Wesołym Miasteczku należy uruchomić dwie nowe kasy, aby zapewnić stabilność systemu. W przeciwnym wypadku system kasowy pozostanie wąskim gardłem obiektu i uniemożliwi wzrost ilości osób odwiedzających.

Maksymalna liczba kas wykorzystywana będzie jedynie w wybrane dni sezonu, najprawdopodobniej jedynie w niedziele i święta w sierpniu oraz przy dalszym wzroście ilości osób odwiedzających również w lipcu.

W przyszłości w dniach, w których nie będzie zakładać się maksymalnego wykorzystania stanowisk kasowych opracowany model optymalizacyjny umożliwi takie planowanie rozpoczęcia pracy poszczególnych stanowisk, które zminimalizuje koszty funkcjonowania systemu.

Literatura

1. Filipowicz B. (1996). Modele stochastyczne w badaniach operacyjnych. Wydawnictwo Naukowo-Techniczne, Warszawa.
2. Filipowicz. B. (1977). Modelowanie i analiza sieci kolejkowych. AGH, Kraków.
3. Gnadenko B.W., Kowalenko I.N. (1971). Wstęp do teorii obsługi masowej. PWN, Warszawa.
4. Kapica M. (2008). Analiza systemu obsługi klienta Śląskiego Wesołego Miasteczka. AE, Katowice.
5. Kleinrock L. (1975). Queueing Systems. Volume I. Theory. Canada.
6. König D., Stoyan D. (1979). Metody teorii masowej obsługi. Wydawnictwo Naukowo-Techniczne, Warszawa.
7. Kopańska-Bródka D. (2006). Wybrane metody badań operacyjnych w zarządzaniu. AE, Katowice.
8. Koźniewska I., Włodarczyk M. (1978). Modele odnowy, niezawodności i masowej obsługi. PWN, Warszawa.
9. Kryński H., Badach A. (1976). Zastosowanie matematyki do podejmowania decyzji ekonomicznych. PWN, Warszawa.
10. Machaczka J. (1990). Modelowanie systemów w organizacji i zarządzaniu. AE, Kraków.
11. Mikuś J. (1993). Metody wspomagania procesu zarządzania, część I. Politechnika Wrocławska, Wrocław.
12. Śląski Park Kultury. Nakładem komitetu budowy WPKiW w Katowicach.
13. Tikhonenko O. (2003). Elementy teorii masowej obsługi. WSP, Częstochowa.

Jerzy Michnik

Akademia Ekonomiczna w Katowicach

ANALIZA KRYTERIÓW WYSTĘPUJĄCYCH W PROCESIE DECYZYJNYM W ZARZĄDZANIU INNOWACJAMI

Wstęp

Wprowadzenie nowego produktu na rynek jest końcowym efektem wielu powiązanych ze sobą działań, które obejmują prace badawczo-rozwojowe, wdrożenie produkcji nowego wyrobu i działania marketingowe. Już na początkowym etapie pojawia się problem wstępnej selekcji projektów. Prowadzone projekty wymagają ciągłego monitorowania oraz decyzji o kontynuacji danego projektu lub jego przerwaniu.

Podczas analizy zarządzania innowacjami już na pierwszy rzut oka zauważa się, że proces ten obejmuje wiele decyzji, w których należy uwzględnić znaczną liczbę różnorodnych kryteriów technicznych, organizacyjnych i rynkowych.

Drugą, istotną cechą charakterystyczną sytuacji decyzyjnych w zarządzaniu innowacjami jest wysoki stopień niepewności. Niepewności te mają swoje źródła wewnątrz przedsiębiorstwa i w jego otoczeniu. Jako wewnętrzne źródła niepewności można wymienić poważne zmiany w technologii lub organizacji oraz znaczne nakłady inwestycyjne, które są niezbędne przy wprowadzaniu innowacji. Głównym źródłem niepewności na zewnątrz przedsiębiorstwa jest niepewna reakcja rynku (klienci, konkurenci) na nowe produkty.

Celem niniejszego opracowania jest analiza procesu wdrażania nowych produktów z punktu widzenia zastosowania metod wielokryterialnych wspomagających podejmowanie decyzji. Głównym punktem zainteresowania są kryteria, ich charakterystyka i wzajemne powiązania. Należy zaznaczyć, że wybór kryteriów nie jest jednoznaczny i zależy od wielu czynników, w szczególności od przyjętych założeń strategicznych przedsiębiorstwa. Wyniki analizy kryteriów mają służyć za podstawę do dalszych badań obejmujących opracowanie odpowiednich metod wielokryterialnych wspomagających podejmowanie decyzji w zarządzaniu nowymi produktami.

1. Metody wielokryterialne w praktyce zarządzania nowymi produktami

Literatura na temat zarządzania innowacjami jest bardzo bogata. Wiele prac dotyczy różnorodnych metod i modeli związanych z podejmowaniem decyzji. Chociaż metody wielokryterialne wydają się przydatnym narzędziem, nie doczekały się one – jak dotychczas – wielu zastosowań w zarządzaniu innowacjami. Nijżej omówiono kilka z nielicznych publikacji, w których przedstawiono zastosowanie metod wielokryterialnych w praktyce.

Zarządzanie projektami BR w Hewlett-Packard.

W procedurze selekcji projektów badawczo-rozwojowych (BR) w powiązaniu ze strategią organizacji, która opracowana została w Hewlett-Packard, wykorzystano metodę AHP [8].

Wstępnie projekty zostały podzielone na kategorie. Wśród wielu możliwości wybrany został dwuwymiarowy podział: poziom zmiany produktu (innowacyjność produktowa) i poziom zmiany procesu (innowacyjność technologiczna) [14]. Na takiej podstawie wyodrębniono trzy kategorie:

1) udoskonalenia, hybrydy i pochodne (niska innowacyjność produktowa i niska innowacyjność technologiczna),

2) nową generację produktów (średnia innowacyjność produktowa i dowolna innowacyjność technologiczna, wysoka innowacyjność technologiczna i średnia lub niska innowacyjność produktowa),

3) produkty przełomowe (wysoka innowacyjność produktowa i dowolna innowacyjność technologiczna).

Następnie przeprowadzano wielostopniowy proces selekcji wstępnej, który miał za zadanie zredukować liczbę projektów do liczby, która miała szansę na realizację przy ograniczonych zasobach. Selekcja ta operowała wieloma sitemi realizującymi wybrane kryteria selekcji, np.

- sitem nr 1: dopasowanie do celów strategicznych
- sitem nr 2: dopasowanie do kompetencji, wielkość rynku, potencjalni partnerzy
- ...
- sitem nr n: dopasowanie do technologii, radykalność, wysyłek marketingowy.

Ostatnim etapem procedury było zastosowanie metody AHP, która posłużyła do porównania projektów osobno w każdej z poszczególnych kategorii. Struktura kryteriów i subkryteriów została wypracowana zespołowo i przedstawiała się następująco (w nawiasach podano wagi przypisane kryteriom; wagi subkryteriów zostały pominięte):

1. Satysfakcja klienta (0,28)
 - 1.1. Poprawa poziomu obsługi.
 - 1.2. Logiczne i dokładne informacje i transakcje.
 - 1.3. Usprawnienie serwisu.
2. Satysfakcja pracowników (0,07)
 - 2.1. Pozytywny wpływ na kontrolę produkcji.
 - 2.2. Poprawa równomierności obłożenia pracą.
 - 2.3. Zwiększenie wydajności lub skuteczności pracowników.
 - 2.4. Poprawa równowagi życie zawodowe/życie prywatne.
 - 2.5. Zwiększenie wiedzy pracowników.
3. Wartość biznesowa (0,46)
 - 3.1. Osiąganie wyników, które są decydujące (krytyczne) w określonym obszarze.
 - 3.2. Minimalizacja ryzyka w okresie wdrożenia i w okresie dojrzałości produktu, poprawa integracji i relacji z partnerami.
 - 3.3. Zapewnienie dodatniego zwrotu z inwestycji (ROI) w okresie 2 lat.
 - 3.4. Zgodność z celami biznesowymi.
4. Skuteczność procesu (0,19)
 - 4.1. Krótki okres wdrażania pracowników.
 - 4.2. Wzmocnienie technologiczne serwisu.
 - 4.3. Redukcja pracy ręcznej i działań nie przynoszących wartości dodanej.
 - 4.4. Zwiększenie samodzielności pracowników.

Ostatnim stadium procedury był wybór projektów z każdej kategorii i ustalenie ostatecznego zestawu projektów, które będą realizowane.

Zastosowanie metody AHP w selekcji nowych produktów

W pracy [2] przedstawiono wykorzystanie metody AHP do selekcji nowych produktów w dużym dziale firmy amerykańskiej. Dział ten jest mocno zaangażowany w rozwój nowych produktów przemysłowych.

Grupa obejmująca przedstawicieli najwyższego kierownictwa działu wzięła udział w dwóch sesjach burzy mózgów. Pierwsza sesja poświęcona była określeniu kryteriów, druga określeniu wag i analizie wrażliwości. W rezultacie uzyskano strukturę AHP, w której rolę kryterium nadrzędnego odgrywał wybór najlepszego projektu nowego produktu.

Procedura porównań parami wyłoniła cztery kryteria, które następnie podzielone zostały na różną ilość subkryteriów. Kryteria (z wagami) oraz subkryteria były następujące:

1. Dopasowanie do podstawowych kompetencji marketingowych (0,285)
 - 1.1. Produkt odpowiada wymaganiom momentu wejścia na rynek wynikającym z docelowych segmentów rynku.

- 1.2. Produkt będzie miał cenę na poziomie (lub poniżej) cen dla docelowych segmentów rynku.
- 1.3. Produkt pasuje do kompetencji logistycznych i marketingowych firmy.
- 1.4. Produkt pasuje do kanałów dystrybucji, w których firma ma silną pozycję.
- 1.5. Produkt pasuje do obecnych linii produkcyjnych.
- 1.6. Produkt pasuje do planów odnoszących się do obszaru sprzedaży, szkoleń i wynagrodzeń.
2. Dopasowanie do podstawowych kompetencji technologicznych (0,227)
 - 2.1. Produkt przynosi klientowi zróżnicowane korzyści i pożytki.
 - 2.2. Tempo produkcji odpowiada popytowi (wystarczająca elastyczność produkcji).
 - 2.3. Produkt jest zaprojektowany według wymagań jakości docelowego segmentu rynku.
 - 2.4. Do produkcji używa się materiałów wysokiej jakości i niskim poziomie zwrotów.
 - 2.5. Produkt pasuje do najlepszej technologii firmy;
 - 2.6. Produkt pozwala na wykorzystanie najlepszych dostawców.
3. Profil zysk-ryzyko w wymiarze finansowym (0,307)
 - 3.1. Całkowite przychody uwzględnione w NPV.
 - 3.2. Całkowite koszty wyrażone w NPV.
4. Ogólna niepewność wyników projektu (0,182)
 - 4.1. Poziom (w %) strat, które nie mogą być określone przez dodatkowe badania.
 - 4.2. Niepewności, które mogą być zredukowane przez zdobycie dodatkowej informacji.

Autorzy w podsumowaniu stwierdzają, że dopasowany do potrzeb konkretnej firmy model AHP może być używany samodzielnie lub zostać włączony w bardziej rozbudowany system wspomagania decyzji.

Selekcja projektów w Advanced Technology Division, Bell Laboratories

Do analizy, rangowania i selekcji projektów BR w Advanced Technology Division (Bell Laboratories) zastosowano model wielokryterialny połączony z graficznym systemem wspomagania decyzji [10]. Autorzy zaproponowali, aby model wielokryterialny został wykorzystany do wstępnej selekcji projektów, które charakteryzują się największym potencjałem. Następnie wśród tych projektów zostaną wybrane te, które zostaną zaakceptowane lub odrzucone bez dalszej analizy. Pozostałe projekty poddane zostaną dalszej analizie za pomocą graficznego systemu wspomagania decyzji (jako przykład podany został, opracowany w firmie Lucent Technologies, Value Creation Model).

Kryteria zostały wybrane przez kierowników działów i kierowników projektów. Pierwsza grupa kryteriów, to kryteria mierzalne (finansowe), które obejmowały dwie pozycje:

1. Kwotę inwestycji.

2. Przewidywany przepływ finansowy w następnych 4 latach, zdyskontowany kosztem kapitału firmy z poprawką uwzględniającą ryzyko. Oszacowanie ryzyka przeprowadzano według przybliżonego, trójwartościowego rozkładu prawdopodobieństwa: wartość pesymistyczna, wartość optymistyczna i wartość najbardziej prawdopodobna.

Druga grupa to kryteria jakościowe dotyczące umiejscowienia potencjalnego projektu w cyklu rozwoju produktu. Przyjęto odwzorowanie czterech jakościowych ocen określenia etapu cyklu życia projektu na skalę liczbową: schyłek (1), rozwinięty (1,5), rozwijający się (2,25), przyszłościowy (1,5)). W tej grupie znalazły się dwa kryteria:

- 1) etap cyklu życia produktu,

- 2) etap cyklu życia własności intelektualnej.

Tak zaprojektowany model zastosowano do zbioru 469 projektów, które zostały podzielone na 3 grupy: 1) złe – odrzucone, 2) bardzo dobre – zaakceptowane, 3) wymagające dalszej analizy. Selekcję przeprowadzono przez subiektywne porównanie projektów pod względem wartości wszystkich 6 kryteriów. W grupie projektów odrzuconych, jak można się domyślać, znalazły się w szczególności projekty zdominowane. Autorzy argumentowali, że znaczne różnice w wartościach ocenianych kryteriów mogą prowadzić do wielu różnych rankingów, zależnie od wybranej metody agregacji kryteriów, dlatego nie zdecydowali się na żadną z nich. Trzecia grupa (wymagające dalszej analizy) była najliczniejsza i do jej analizy użyto jakościowej metody Value Creation Model (VCM), która była używana w Advanced Technology Division od kilku lat. Autorzy wyrazili opinię, że w miarę jak menedżerowie przyzwyczajają się do modelu wielokryterialnego i obdarzą go większym zaufaniem, większa część projektów będzie klasyfikowana za jego pomocą – np. za pomocą rankingów – do pierwszych dwóch kategorii, a mniejsza część projektów będzie wymagała dalszej, czasochłonnej analizy.

MCDM w międzynarodowej firmie produkcyjnej

W artykule M.S. Morcosa [11] zaproponowany został wielokryterialny model oparty na metodzie SMART (Simple Multi-Attribute Rating Technique) opracowanej w latach 70. ubiegłego wieku [6; 7].

Projekty zostały podzielone pod względem wielkości na cztery kategorie: małe, średnie, duże i bardzo duże. Redukcję ilości wariantów decyzyjnych uzyskano definiując 36 pakietów strategicznych (odpowiednio 2, 3, 3, 2 warianty dla każdej kategorii projektów).

Ocena projektów oparta została na konfrontacji wykorzystania zasobów i potencjalnych korzyści. Za najważniejszy zasób uznano pieniądź, co ograniczyło analizę zasobów do jednego kryterium – kosztu inwestycji.

Korzyści zostały podzielone na materialne i niematerialne. Za korzyści materialne przyjęto utrzymanie zyskowności w krótkim horyzoncie czasowym. Korzyści niematerialne obejmowały dwie pozycje: niezawodność (zwiększanie niezawodności projektów) i ryzyko (minimalizowanie ryzyka związanego z realizowanym portfelem projektów BR).

W kolejnym kroku oszacowano koszty i korzyści dla każdego wariantu. Koszty wariantu i jego potencjalne zyski (korzyści materialne) wyrażono w milionach funtów brytyjskich, natomiast korzyści niematerialne określano w skali 0-100 (dla ryzyka skala była odwrócona – najniższemu ryzyku przypisano 100 pkt., a najwyższemu 0 pkt.).

Aby uzyskać możliwość wizualizacji wariantów na dwuwymiarowych wykresach (na wzór wykresów zwrot-ryzyko w modelach Markowitza), dokonano ujednolicenia skali oraz agregacji kryteriów dotyczących korzyści. Zyski zostały przeliczone na skalę 0-100 według schematu proporcjonalnego (wariant o najwyższym zyskach otrzymał 100 pkt., a wariant o najniższych – 0 pkt.; pozostałe otrzymały wartości pośrednie). Następnie oszacowano wagi poszczególnych kryteriów korzyści i za pomocą modelu addytywnego obliczono ważone oceny dla wszystkich wariantów strategicznych.

W efekcie uzyskano model dwukryterialny, w którym jedno kryterium stanowiły koszty inwestycji, a drugie zagregowana miara korzyści. W ramach tego modelu wyselekcjonowano projekty niezdominowane, z których następnie wybrano rozwiązanie końcowe.

2. Analiza kryteriów w zarządzaniu innowacjami

Analiza literatury na temat zarządzania innowacjami wskazuje, że już na poziomie werbalizacji celu głównego, który ma wskazywać kierunek dążeń w procesie wyboru nowych produktów, sytuacja nie jest jednoznaczna. Jako możliwe cele główne proponuje się:

- skuteczne funkcjonowanie [3],
- wyniki finansowe produktu [4],
- rozwój przedsiębiorstwa [1, s. 264],
- wzmocnienie (lub przynajmniej utrzymanie) pozycji konkurencyjnej [9, s. 40],
- funkcjonowanie, dokonania (performance) [12],
- lepszą, wyróżniającą pozycję na rynku [13, s. xxxvii],
- budowanie trwałej przewagi konkurencyjnej [5, s. 100].

Jak wynika z powyższego zestawienia, cel główny bywa formułowany na różnym poziomie ogólności i może być nakierowany na różne obszary aktywności przedsiębiorstwa.

Cel główny, w mniejszym lub większym stopniu, implikuje podstawowe kryteria, które pojawiają się w dyskusji na temat zarządzania innowacjami.

Szeroki zestaw kryteriów obejmujący wiele aspektów działalności przedsiębiorstwa, które mogłyby posłużyć za kanwę modelu wielokryterialnego przedstawił J. Baruk (w oryginale występują one jako cele częściowe, prowadzące do celu głównego, którym jest rozwój przedsiębiorstwa) [1, s. 264]:

- dostosowanie wielkości i struktury produkcji do zmieniających się potrzeb rynku,
- wzrost jakości,
- optymalizacja wykorzystania materiałów, paliw, itp.,
- doskonalenie środków produkcji.

Przedostatnie kryterium można uogólnić, mówiąc o optymalizacji zasobów, uwzględniając również wykorzystanie maszyn i urządzeń oraz zasobów ludzkich.

Analiza istotnych aspektów zarządzania portfelem nowych produktów doprowadziła autorów pracy [3] do sformułowania, jak to określili „strategicznych kryteriów w ocenie portfela nowych produktów”:

- zgodność ze strategią firmy,
- równowaga portfela,
- właściwy rozmiar portfela.

Przez właściwy rozmiar portfela rozumie się ograniczoną liczbę projektów uwzględniającą zasoby firmy tak, aby nie następowały opóźnienia w realizacji projektów, wynikające z konkurencji o rzadkie zasoby.

Równowaga portfela obejmuje z kolei szereg subkryteriów odnoszących się do równowagi pomiędzy:

- projektami długo- i krótkoterminowymi (czynnik czasu),
- projektami o niskim i wysokim ryzyku,

Portfel jest również w równowadze, gdy obejmuje projekty zróżnicowane technologicznie i nakierowane na zróżnicowane rynki.

Badając innowacyjność z perspektywy przedsiębiorstwa, Danneels i Kleinschmidt przeprowadzili statystyczną analizę wieloczynnikową danych empirycznych z przedsiębiorstw kanadyjskich [4]. Z badań tych wynika, że duże znaczenie w ocenie (ex post) nowych produktów odgrywają wyniki finansowe produktu. Na tak określony czynnik (factor) złożyły się następujące kryteria:

- ocena zyskowności (względem akceptowalnego minimum),
- ocena sprzedaży (względem innych produktów),
- zyski (względem innych produktów),
- stopień realizacji planowanej (docelowej) sprzedaży,
- stopień realizacji planowanego (docelowego) zysku.

Kryteria te mogą być podstawą modelu, w którym nacisk położony jest na wyniki finansowe przedsiębiorstwa, uzyskane w konsekwencji wprowadzania innowacji.

2.1. Próba klasyfikacji kryteriów

Złożoność problemów decyzyjnych, które występują w zarządzaniu innowacjami powoduje, że podczas analizy pojawia się znaczna ilość kryteriów decyzyjnych. Kryteria te mają bardzo różnorodny charakter. W pierwszym rzędzie pojawiają się one na różnych poziomach ogólności, a więc zajmą różne miejsce w hierarchii kryteriów: jako kryteria lub subkryteria.

Zróznicowany charakter kryteriów wynika z również z tego, że odnoszą się one do różnych obszarów aktywności przedsiębiorstwa. W takiej sytuacji próba klasyfikacji kryteriów może znacząco pomóc w analizie problemu i budowie odpowiedniego modelu decyzyjnego.

Jednym z podstawowych paradygmatów teorii zarządzania jest dychotomiczny podział na wnętrze i otoczenie zewnętrzne firmy. Wychodząc z tego punktu widzenia można próbować analogicznego podziału na kryteria wewnętrzne i zewnętrzne. Do kryteriów, które można uznać za wewnętrzne należą np. satysfakcja pracowników, wzrost jakości, optymalizacja wykorzystania zasobów, zgodność ze strategią firmy. Do kryteriów o charakterze zewnętrznym można zaliczyć satysfakcję klienta, cykl życia produktu, cykl życia własności intelektualnej. Jednakże większość kryteriów, szczególnie finansowych, nie daje się jednoznacznie zaklasyfikować do jednej z powyższych grup. Są one uwarunkowane zarówno czynnikami wewnętrznymi firmy, jak i sytuacją zewnętrzną. Na przykład koszty inwestycji w dany projekt zależą od sytuacji finansowej firmy, jej zasobów gotówkowych, ale także od sytuacji rynkowej, w tym cen (koszty zakupu odpowiednich urządzeń i materiałów) i stóp procentowych (finansowanie projektu środkami zewnętrznymi).

Bardziej przydatna w analizie może być klasyfikacja oparta na merytorycznej treści kryteriów. Analiza przykładów praktycznych oraz literatury teoretycznej sugeruje wyodrębnienie następujących grup kryteriów:

- strategiczne, oceniają warianty decyzyjne z punktu widzenia planów strategicznych i długoterminowych celów,
- organizacyjno-biznesowe, powiązane z umiejętnościami organizacyjnymi, kompetencjami marketingowymi, logistycznymi itp.,

- techniczne, obejmują kwestie związane z techniczną stroną przygotowania nowych produktów (potencjał badawczo-rozwojowy, kompetencje technologiczne),
- finansowe, dotyczą kosztów inwestycji w przygotowanie nowych produktów oraz spodziewanych przychodów ze sprzedaży nowych produktów,
- rynkowe, związane z szeroko rozumianymi czynnikami rynkowymi, w tym reakcją klientów i konkurencji na nowy produkt.

2.2. Zależności pomiędzy kryteriami

W metodach wielokryterialnych wspomaganie decyzji opiera się najczęściej na procedurze agregacji ocen częściowych, dotyczących poszczególnych kryteriów. Zagregowana ocena wariantów decyzyjnych ma najprostsza formę, gdy nie trzeba uwzględniać zależności pomiędzy ocenami. Stąd częstym rozwiązaniem jest przyjęcie założenia o niezależności kryteriów.

Jednakże w wielu problemach praktycznych kryteria mogą być na tyle mocno powiązane, że założenie o niezależności staje się zbyt uproszczeniem i może prowadzić do gorszych rezultatów niż w modelu uwzględniającym powiązania pomiędzy kryteriami.

We wszystkich przykładach metod wielokryterialnych, przytoczonych w tej pracy, wykorzystane modele wielokryterialne zakładają niezależność kryteriów. Jednakże już pobieżna analiza wskazuje, że może to być zbyt uproszczenie. Na przykład w modelu AHP opisanym w pracy [2] ocena niepewności wyników projektu (kryterium 4) jest zależna od dopasowania do kompetencji marketingowych (kryterium 1.) i kompetencji technologicznych (kryterium 2). Dodatkowo należałoby się spodziewać zależności oceny przychodów i kosztów (kryterium 3) od niepewności wyników. Podobna sytuacja występuje w modelu opisanym przez Morcosa [11], gdzie zależności mogą wystąpić pomiędzy kryteriami materialnymi (przychody i koszty) i kryteriami niematerialnymi (niezawodność i ryzyko).

Z kolei w modelu opracowanym w Advance Technology Division (Bell Lab.) [10] jakościowe kryteria oceny dotyczące cyklu życia produktu mogą mieć stosunkowo silny wpływ na ocenę rozkładu przepływów finansowych związanych z produktem.

Pomiędzy kryteriami, proponowanymi w opracowaniach teoretycznych, też można znaleźć przykłady powiązań, mogących znacząco wpłynąć na wyniki modelu wielokryterialnego. Jako przykład dwustronnej zależności można wskazać, wymienione w monografii [1], doskonalenie środków produkcji i wzrost jakości, dwa kryteria, które mogą wpływać na siebie wzajemnie.

Podsumowanie

Zarządzanie innowacjami jest przykładem złożonego procesu, który ma istotne znaczenie dla wielu firm. W procesie tym występuje duża liczba różnorodnych kryteriów odnoszących się do różnych aspektów sytuacji i działalności firmy. Kryteria te różnią się znaczeniem merytorycznym oraz własnościami formalnymi. W szczególności, agregacja kryteriów o różnym charakterze wymaga określenia odpowiednich skal, pozwalających na łączenie ocen kryteriów ilościowych z jakościowymi. Ze względu na dużą ilość kryteriów pomocne mogą być agregacje cząstkowe (przykład takiej procedury w celu wizualizacji wyników na płaszczyźnie przedstawiono w artykule [11]).

Przegląd literatury wskazuje na rosnące zainteresowanie usprawnianiem procesu zarządzania innowacjami. Nieliczne przykłady zastosowań metod wielokryterialnych oparte są na modelach, w których niepewność traktowana jest w sposób uproszczony, a zależności między kryteriami w ogóle nie są uwzględniane. Opisane próby zastosowania metod wielokryterialnych opierają się na wieloetapowych procedurach, w których, zwykle ad hoc, łączone są różne metody i techniki (np. często jako składnik tych procedur występuje metoda AHP).

Nasuwa się oczywisty wniosek, że istnieje szerokie pole zastosowań dla zaawansowanych metod wielokryterialnych, które w większym stopniu pozwolą na usprawnienie procesu zarządzania innowacjami. Za najważniejsze postulaty dalszych badań należy uznać:

- określenie w miarę pełnego zestawu kryteriów istotnych w zarządzaniu innowacjami,
- wykorzystanie metod uwzględniających zależności pomiędzy kryteriami,
- wprowadzenie bardziej zaawansowanych technik uwzględnienia niepewności w modelach wielokryterialnych.

Literatura

1. Baruk J. (2006). Zarządzanie wiedzą i innowacjami. Adam Marszałek, Toruń.
2. Calantone R.J., Benedetto C.A., Schmidt J.B. (1999). Using the Analytic Hierarchy Process in new Product Screening. An International Publication of The Product Development & Management Association, 16 (1), 65-76.
3. Cooper R.G., Edgett S.J., Kleinschmidt E.J. (1999). New Product Portfolio Management: Practices and Performance. Journal of Product Innovation Management, 16 (4), 333-351.

4. Danneels E., Kleinschmidt E.J. (2001). Product Innovativeness from the Firm's Perspective: its Dimensions and their Relation with Project Selection and Performance. *Journal of Product Innovation Management*, 18 (6), 357-373.
5. Dodgson M., Gann D., Salter A. (2008). *The Management of Technological Innovation*. Oxford University Press, Oxford.
6. Edwards W. (1971). Social Utilities. *Engineering Economist* (Summer Symposium Series), 6.
7. Edwards W. (1977). How to Use Multi-attribute Utility Measurement for Social Decision-making. *IEEE Transactions on Systems, Man and Cybernetics*, 7 (5), 326-40.
8. Englund R.L., Graham R.J. (1999). From Experience: Linking Projects to Strategy. *Journal of Product Innovation Management*, 16 (1), 52-64.
9. Jasiński A.H. (2006). *Innowacje i transfer techniki w procesie transformacji*. Centrum Doradztwa i Informacji, Difin, Warszawa.
10. Linton J.D. et al. (2000). Selection of R&D Projects in a Portfolio. W: *Engineering Management Society, 2000. Proceedings of the 2000 IEEE*, 506-511.
11. Morcos M.S. (2008). Modelling Resource Allocation of R&D Project Portfolios Using a Multi-criteria Decision-making Methodology. *International Journal of Quality and Reliability Management*, 25 (1), 72.
12. Siguaw J.A., Simpson P.M., Enz C.A. (2006). Conceptualizing Innovation Orientation: A Framework for Study and Integration of Innovation Research. *Journal of Product Innovation Management*, 23 (6), 556-574.
13. Westland J.C. (2008). *Global Innovation Management*. Palgrave, Macmillan.
14. Wheelwright S., Clark K. (1992). Creating Project Plans to Focus Product Development. *Harvard Business Review* (March-April).

Bogusław Nowak

Akademia Ekonomiczna w Katowicach

PLANOWANIE INWESTYCJI BUDOWLANEJ – ANALIZA SYMULACYJNA

Wprowadzenie

Skuteczne planowanie inwestycji ściśle związane jest z koniecznością podejmowania wielu trafnych decyzji w przypadku dostatecznej wiedzy dotyczącej konsekwencji dokonywanych wyborów. Celem zminimalizowania ryzyka braku dostatecznej informacji korzystne wydaje się wykorzystanie technik komputerowych umożliwiających dokonanie symulacji efektów, do których prowadzi podjęcie określonych decyzji [3]. Przeprowadzane za pomocą programu komputerowego symulacje pozwalają na uwzględnianie niepewności, jak i dynamiki analizowanych przypadków [4]. W niniejszej pracy wykorzystano najbardziej rozpowszechnione narzędzie, jakie stanowi arkusz kalkulacyjny Excel. Wspomniany arkusz kalkulacyjny potraktowany został jako narzędzie umożliwiające pozyskanie wiedzy niemożliwej do zdobycia w trakcie rzeczywistej realizacji projektu. Zaprezentowany przykładowy projekt* wymusił na przedsiębiorstwie konieczność jego podziału na kilka istotnych etapów. Wymóg ten spowodowany został dużą złożonością zadania, a zaproponowany podział umożliwił łatwiejszą jego kontrolę. W niniejszej pracy skupiono się na jednym tylko etapie dotyczącym wyboru projektu oraz sposobu jego realizacji.

Celem pracy jest przedstawienie możliwości wykorzystania drzewa decyzyjnego i symulacji komputerowej do wspomagania planowania projektu inwestycyjno-budowlanego**. W analizowanym zagadnieniu ze względu na jego cel przyjęto wiele uproszczeń, pomijając kwestie nieistotne z punktu widzenia niniejszej pracy. Należała do nich np. analiza bezpieczeństwa nowego rynku (analizowany projekt realizowany był poza granicami Polski). Nie

* Termin projekt określany przez współczesną literaturę i posiadający wiele definicji określany według jego szczególnych zastosowań stanowi dla przykładu tymczasowe przedsięwzięcie podejmowane w celu wytworzenia unikalnego wyrobu lub usługi.

** Jeden z 4 typów projektów realizowanych przez BT ZWUS. Projekty te realizowane metodą PMBoK stanowią największą grupę przedsięwzięć realizowanych przez przedsiębiorstwo.

uwzględniono również faktu, że branża, do której należy przedsiębiorstwo*, zatem tak jego zasoby, jak i oferowane wyroby czy usługi wymagają uzyskania specjalnych certyfikatów i dopuszczeń. Na potrzeby niniejszego opracowania oprócz drzewa decyzyjnego przygotowany został w arkuszu Excel model symulacyjny, dzięki któremu oszacowano prawdopodobieństwo uzyskania założonej marży projektu.

1. Projekty inwestycyjno-budowlane

Rozważana w pracy grupa projektów należąca do największych z pośród realizowanych przez przedsiębiorstwa przedsięwzięć składa się z 2 wzajemnie powiązanych części:

1) planowania realizowanego przez przedsiębiorstwo według np. korporacyjnych założeń składającego się z:

- oceny możliwości i ryzyka,
- kalkulacji tzw. planu finansowego,
- oferty stanowiącej produkt końcowy etapu planowania,

2) realizacji przedsięwzięcia również realizowanego według wewnętrznych ustaleń (planu projektu) składającego się z głównych etapów:

- produkcji urządzeń na potrzeby projektu,
- adaptacji inżynierskich niezbędnych celem przystosowania własnych produktów do istniejącego systemu klienta,
- właściwej realizacji obiektu (etapu budowy),
- opieki gwarancyjnej zainstalowanych urządzeń u klienta.

W analizowanej grupie projektów w początkowym etapie procesu decyzyjnego istnieje przeważnie konieczność dokonania wyboru pomiędzy możliwymi do realizacji projektami tak krajowymi, jak i zagranicznymi związanymi głównie z pozyskaniem nowych rynków. Rozpatrując udział w przetargu należy uwzględnić możliwość spełnienia przez przedsiębiorstwo warunków przetargowych oraz reżimu czasu, w jakim wymagana jest realizacja zadania. W przypadku pozyskania kontraktu na zagranicznym rynku należy uwzględnić możliwość realizacji projektu przy współpracy z obcą, ale doświadczoną w tego typu przedsięwzięciach firmą. Inne występujące tu możliwości to realizacja projektu z lokalnym przedstawicielem lub ograniczenie się do roli jedynie dostawcy urządzeń. Przyjmowane wartości kosztów poszczególnych opcji wynikają przede wszystkim z zakresu przyjętych zadań. W przypadku, gdy firma

* Rozpatrywane przedsiębiorstwo scharakteryzować można jako organizację o projektowym stylu zarządzania. Działalność jego dotyczy specyficznej gałęzi przemysłu wyróżniającej się szczególną dbałością o bezpieczeństwo oferowanych urządzeń oraz usług. Ze względu na szczególny charakter jego działalności realizacja każdego projektu wymaga wyjątkowego przygotowania i dbałości w realizacji, jak również fachowości kierowników projektów.

pełni rolę generalnego wykonawcy, najistotniejszym składnikiem kosztów są wydatki związane z organizacją całości przedsięwzięcia. W przypadku, gdy przedsiębiorstwo wchodzi w skład konsorcjum, koszty te znacząco maleją tak, by przy roli wyłącznie dostawcy urządzeń spaść zaledwie do kilkunastu procent ich pierwotnej wartości.

Do innych kosztów istotnie wpływających na wartość przedsięwzięcia należą koszty wykonawców. Koszty te w przypadku własnej grupy wykonawców mogą znacząco przewyższać koszty zewnętrznych wykonawców pozyskiwanych w ramach najkorzystniejszej oferty. Na wysokie koszty własnych pracowników wpływają zasadniczo takie czynniki, jak koszty administracyjne, koszty korporacyjne, ale także koszty utraconych korzyści wynikających z zaangażowania na określony czas grupy własnych pracowników i braku możliwości ich wykorzystania w innych równolegle prowadzonych przedsięwzięciach np. ze względu na odległość budowy od innych równolegle realizowanych zadań. Poza tym kompetencje i wiedza własnych pracowników są często wyższe od tych, które wymagane są przy realizacji rozważanego projektu. Ta wszechstronność również jest przyczyną wzrostu ponoszonych przez przedsiębiorstwo kosztów [1].

Do kosztów dodatkowo wpływających na całkowity bilans przedsięwzięcia należą również koszty produkcji. W przypadku wyboru własnych urządzeń istnieje możliwość dostosowania tych kosztów do ceny oferty ostatecznie kreowanej przez dział sprzedaży. Odpowiednia zmiana marży powoduje w efekcie zmianę oferty końcowej zależnej od stawianych przedsięwzięciu priorytetów, dlatego też w przypadku włączenia do oferty rozwiązań konkurencyjnych nie istnieje możliwość tak elastycznego dopasowania oferty cenowej do przypuszczalnych oczekiwań klienta.

2. Drzewo decyzyjne i symulacja komputerowa

Drzewo decyzyjne to graficzna metoda wspomagania procesu decyzyjnego, pozwalająca na wyznaczenie najkorzystniejszego rozwiązania z punktu widzenia rozważanego kryterium. Metoda ta wymaga od decydenta znajomości prawdopodobieństw zaistnienia stanów natury w istotny sposób wpływających na efekty podejmowanych decyzji. Zapisanie procesu decyzyjnego w postaci grafu pozwala w łatwy sposób dostrzec możliwe zagrożenia będące następstwem przyjętych działań, ale pozwala również z wyprzedzeniem na nie reagować. Budowa drzewa wymusza na decydencie zapisanie w logiczny sposób struktury całego zadania. Istotne jest zachowanie odpowiedniej kolejności wypełniania poszczególnych gałęzi drzewa, poczynając od tzw. korzenia poprzez określone decyzje, możliwe stany natury aż po stany końcowe, w których

proces może się znaleźć w wyniku zastosowania ustalonej strategii postępowania i zajścia określonych stanów natury. Następnie podstawiając odpowiednie wartości i dokonując prostych obliczeń w łatwy sposób możliwe staje się wskazanie rozwiązania, które jest najkorzystniejsze z punktu widzenia analizowanego kryterium [6].

Do wyboru najkorzystniejszej ścieżki przebiegu procesu wykorzystać możemy np. regułę maksymalizacji oczekiwanej korzyści. Reguła ta wykorzystując rozkład prawdopodobieństwa stanów natury pozwala obliczyć oczekiwane korzyści dla każdej z rozpatrywanych decyzji. Ostatecznie wybór pada na tę decyzję, dla której wartość oczekiwanej korzyści jest największa [5].

Informacja, jaką uzyskujemy korzystając z drzewa decyzyjnego jest dość uboga. Rozwiązując problem koncentrujemy się bowiem na wyznaczeniu takiej strategii postępowania, przy której oczekiwana korzyść jest maksymalna. Abstrahujemy zatem od ryzyka, jakie wiąże się z wyborem takiego, a nie innego rozwiązania. Jeżeli korzyści realizowane w węzłach końcowych drzewa są podane w postaci liczb rzeczywistych, to na podstawie struktury drzewa jesteśmy w stanie ustalić rozkład prawdopodobieństwa korzyści uzyskiwanej w wyniku zastosowania określonej strategii postępowania. Uzyskanie takiego rozkładu pozwala na oszacowanie ryzyka, rozumianego np. jako prawdopodobieństwo uzyskania korzyści na poziomie niższym od wymaganego. W przypadku drzewa o skomplikowanej strukturze wygenerowanie rozkładu w opisany powyżej sposób może być jednak dość kłopotliwe. Metoda ta zawodzi natomiast, gdy korzyści realizowane w węzłach końcowych nie są wyrażone za pomocą wartości zdeterminowanych, lecz opisane rozkładami prawdopodobieństwa. W takim wypadku skorzystać natomiast możemy z metod symulacyjnych. Symulacja jako narzędzie badawcze stanowi źródło wiedzy niemożliwej do uzyskania w trakcie rzeczywistych obserwacji analizowanych zjawisk. W pierwszej kolejności budujemy model analizowanego zjawiska, a następnie przeprowadzamy na nim ciąg eksperymentów symulujących rzeczywiste otoczenie. Eksperymenty te realizowane są najczęściej za pomocą programów komputerowych. Niejednokrotnie wykorzystanie oprogramowania komputerowego stanowi jedyną możliwość pozyskania interesującej nas wiedzy. Rozwiązane analizowanego problemu opisanego modelem matematycznym może być uzyskane zarówno za pomocą metod analitycznych, jak i numerycznych. W niniejszym opracowaniu wykorzystano podejście numeryczne. Wykorzystany model symulacyjny opisuje relacje pomiędzy poszczególnymi składnikami analizowanego zjawiska oraz dostarcza informacji pozwalających na ocenę proponowanego rozwiązania.

Wejście modelu symulacyjnego stanowią zmienne decyzyjne, których wartości określa decydent oraz zmienne stanu. Na wyjściu modelu otrzymujemy zmienne wyjściowe wykorzystywane jako miernik efektywności analizowanego

rozwiązania. Metoda proponowana w niniejszej pracy stanowi połączenie drzewa decyzyjnego oraz symulacji komputerowej. Zakładamy, że korzyści realizowane w węzłach końcowych drzewa decyzyjnego opisane są rozkładami trójkątnymi. Interesować nas będzie wyznaczenie rozwiązania gwarantującego uzyskanie maksymalnej oczekiwanej korzyści, a następnie przeanalizowanie ryzyka związanego z jego realizacją. W tym celu skorzystamy z dwuetapowej procedury. W kroku pierwszym wykorzystane zostanie drzewo decyzyjne. W etapie drugim dla tak wyznaczonego rozwiązania przeprowadzone zostaną eksperymenty symulacyjne. W ich efekcie uzyskana zostanie dodatkowa informacji na temat konsekwencji, do jakich prowadzić może zastosowanie rozwiązania wyznaczonego w etapie pierwszym.

3. Opis analizowanego zagadnienia

Analizowany w niniejszym opracowaniu problem jest przykładem wieloetapowego podejmowania decyzji w projektach należących do grupy inwestycyjno-budowlanych. Przyjęty w przykładzie rozkład trójkątny ostatecznej ceny oferty wynika z możliwych do wyboru opisanych wyżej opcji, jak np. wybór własnych instalatorów lub instalatorów z firmy współpracującej w przedsięwzięciu. W rozważanym przykładzie mamy do czynienia z wieloma rodzajami ryzyka wynikającymi z wejścia na nowy, dotąd nieznaną rynek. Na tym etapie realizacji projektu przed decydem stoi wiele koniecznych do podjęcia decyzji. Wygodnym narzędziem w określaniu strategii postępowania okazało się zastosowanie drzewa decyzyjnego. Z uwagi na fakt, że rozwiązując problem oparto się na regule maksymalizacji wartości oczekiwanej, wykorzystując to narzędzie uzyskano w efekcie informację o oczekiwanej marży. Celem pozyskania dodatkowych informacji np. takich jak prawdopodobieństwo, że wynik będzie gorszy niż próg określony na wstępie zastosowano symulację komputerową do wyznaczenia rozkładu marży uzyskiwanej w efekcie zastosowanej strategii. Szacując rozkłady prawdopodobieństwa korzyści uzyskiwanych w węzłach końcowych drzewa wykorzystano zarówno bazę wiedzy zawierającą informacje na temat wcześniej realizowanych przedsięwzięć, jak też informacje pozyskane od ekspertów mających doświadczenie w realizacji podobnego typu projektów.

4. Przykład liczbowy

Jak wspomniano wyżej, analizowane zagadnienie ze względu na swoją objętość podzielone zostało na kilka istotnych etapów. W niniejszej pracy skupiono się wyłącznie na etapie planowania łącznie z decyzjami organizacyjnymi dotyczącymi przyszłego przedsięwzięcia.

W pierwszym etapie rozpatrywano następujące decyzje:

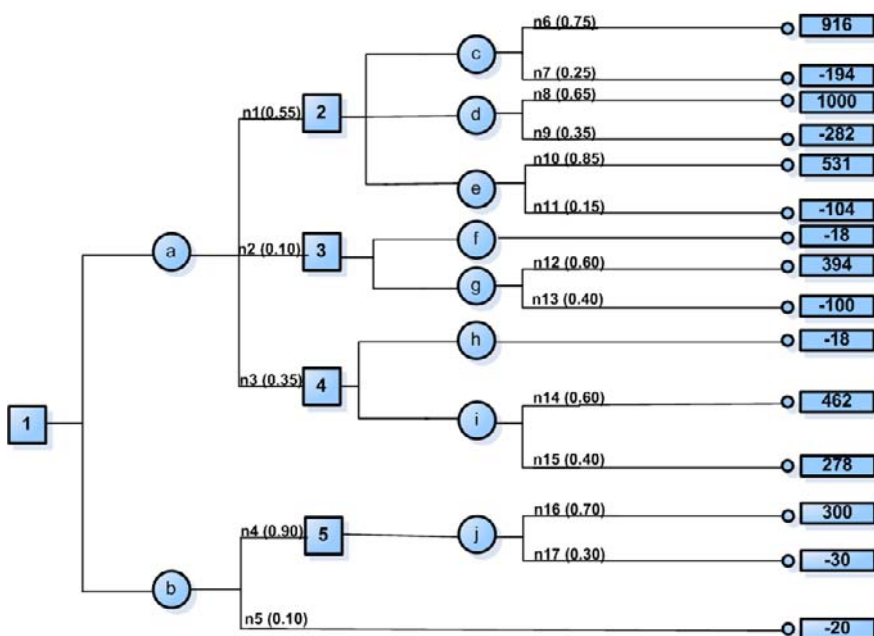
1) realizację projektu na rynku krajowym lub realizację projektu na rynku zagranicznym.

Odpowiedzialne za wybór jednej z powyższych opcji było kierownictwo przedsiębiorstwa. Decyzję o wyborze opcji drugiej podjęło na podstawie raportu z analizy rynku i możliwości realizacyjnych przedsiębiorstwa dokonanych przez własnych specjalistów.

W drugim etapie rozpatrywano kolejne decyzje:

- 2) realizację projektu przy współpracy z lokalną (miejscową) firmą,
- 3) realizację projektu samodzielnie jako generalny wykonawca,
- 4) realizację projektu w konsorcjum z lokalną firmą,
- 5) realizację projektu w funkcji wyłącznie dostawcy urządzeń.

Odpowiedzialny za wybór powyższych w przeciwieństwie do decyzji strategicznych był kierownik projektu przy wsparciu powołanej grupy projektowej składającej się z przedstawicieli przedsiębiorstwa. Przykład drzewa decyzyjnego przedstawia rys. 1.



Rys. 1. Fragment drzewa decyzyjnego opisującego wybór projektu oraz wybór sposobów jego realizacji

Wartości końcowe przedsięwzięcia uzyskano korzystając ze standardowego arkusza kalkulacyjnego projektu* szeroko stosowanego przez przedsiębiorstwo rozbudowując go o dwie dodatkowe opcje: optymistyczną oraz pesymistyczną. Arkusz ten umożliwił uwzględnienie wszelkich narzutów oraz kosztów administracyjnych i korporacyjnych, jak również kosztów związanych z posiadanymi licencjami. Końcowa wartość kosztów powiększona o wnioskowaną marżę stanowiła ostateczną wartość oferty. Opcja najbardziej prawdopodobna stanowiła rzeczywistą wartość projektu oszacowaną przez przedsiębiorstwo podczas etapu kalkulacji projektu. Opcje optymistyczną oraz pesymistyczną, o których mowa wyżej uzyskano w następujący sposób:

1. W przypadku generalnego wykonawstwa w opcji optymistycznej przyjęte koszty ustalono na podstawie wyłącznie własnej produkcji oraz dzięki wykonawcom pochodzącym z firm zewnętrznych, a ponadto lokalnych, którzy po krótkim szkoleniu będą w stanie sami sprawnie dokonywać wszelkich instalacji urządzeń oferowanych w ramach rozważanego projektu. Przyjęto również brak jakichkolwiek problemów z powiązaniem nowych urządzeń z systemem klienta, gdyż kontrakt obejmuje wyłącznie zabudowę i uruchomienie urządzeń wykonawczych z powiązaniem do istniejącego systemu nadrzędnego, sterującego klienta. Nie uwzględniono ponadto kwestii kosztów magazynowania wyprodukowanych już urządzeń. Przyjęto dostawę na miejsce i natychmiastową ich zabudowę.

2. Opcja pesymistyczna zakłada natomiast wystąpienie wielu problemów, takich jak konieczność częściowego zakupu konkurencyjnych urządzeń ze względu na brak możliwości realizacji zlecenia przez własną produkcję wyłącznie na potrzeby analizowanego przedsięwzięcia. Przyjęto ponadto realizację instalacji wyłącznie przez własnych pracowników, np. ze względu na brak zainteresowania lokalnych firm.

Dokonując szczegółowej analizy wybranego projektu w pierwszym kroku (węzeł decyzyjny 1) kierownictwo przedsiębiorstwa musi dokonać wyboru pomiędzy dwoma decyzjami:

- a) realizacją średniego projektu na nowym – nieznanym rynku,
- b) realizacją małego projektu na znanym rynku krajowym.

Kolejnymi węzłami drzewa decyzyjnego są węzły losowe a oraz b. Z węzła losowego a wychodzą trzy krawędzie odpowiadające realizacji występujących stanów natury:

- n1 przedsiębiorstwo jest w stanie spełnić warunki przetargowe (prawdopodobieństwo 0,55),
- n2 przedsiębiorstwo nie spełnia warunków przetargowych (prawdopodobieństwo 0,1),

* Arkusz Excel z uwzględnionymi możliwymi składnikami kosztowymi projektu, tj. materiałami, robocizną, transportem, ryzykiem, administracją, gwarancją itp.

n3 przedsiębiorstwo spełnia warunki przetargowe, jednak wymagany czas realizacji projektu jest krótszy od możliwości realizacyjnych firmy (prawdopodobieństwo 0,35).

Z kolei z węzła losowego b wychodzą tylko dwie* krawędzie odpowiadające realizacji stanów natury:

n4 Przedsiębiorstwo spełnia warunki przetargowe (prawdopodobieństwo 0,9),

n5 Przedsiębiorstwo nie spełnia warunków przetargowych (prawdopodobieństwo 0,1).

W wyniku realizacji stanów natury n1, n2, n3, n4 pojawiają się węzły decyzyjne 2, 3, 4, 5, w których możliwe do podjęcia są następujące decyzje:

Węzeł decyzyjny 2:

c. Realizacja projektu przy współpracy z lokalną (miejscową) firmą.

d. Realizacja projektu samodzielnie jako generalny wykonawca.

e. Realizacja projektu w konsorcjum z lokalną firmą.

Węzeł decyzyjny 3:

f. Wycofanie z przetargu.

g. Nawiązanie współpracy z innym startującym w przetargu przedsiębiorstwem tak, by stać się wyłącznie dostawcą urządzeń.

Węzeł decyzyjny 4:

h. Wycofanie z przetargu.

i. Negocjacje z działem produkcji oraz kooperantami celem startu w przetargu jako dostawca urządzeń i nadania zadaniu wyższego priorytetu kosztem możliwych opóźnień innych realizowanych aktualnie projektów.

Węzeł decyzyjny 5:

j. Realizacja projektu samodzielnie jako generalny wykonawca.

Ostatni krok stanowi identyfikacja prawdopodobnych stanów natury od n5 do n17, a w ich wyniku możliwych wypłat (bez uwzględnienia wypłat w kolejnych etapach).

Stan n5, podobnie jak n7, n9, n11, n13, n17, oznacza rezygnację lub niepowodzenie przyjętych decyzji. Dla przykładu, n6 oznacza 75% prawdopodobieństwo sukcesu przyjętej strategii nawiązania współpracy z lokalną, miejscową firmą. Sukces ten dotyczy porozumienia co do podziału pracy, obowiązków i zysków. Jak brak skutkuje wycofaniem z przetargu i poniesieniem przez przedsiębiorstwo z tego tytułu z 25% prawdopodobieństwem oznaczonym jako n7 strat w wysokości 194 tys. zł. Podobnie stan natury n8 oznacza sukces podjętych decyzji obliczonych jako przychód w wysokości 1000 tys. zł. Wycofanie z przetargu spowodowane powstałymi problemami pociągnie za sobą ujemny bilans działalności wynoszący 282 tys. zł straty oznaczony jako n9. Podobnie wycofanie z przetargu spowodowane niespełnieniem warunków przetargowych

* Na rynku krajowym w branży, w której działa przedsiębiorstwo to klient zawsze określa czas realizacji przedsięwzięcia. Oferenci oszacowując własne możliwości realizacyjne muszą uwzględnić ten postawiony na wstępie wymóg.

(decyzja f) lub niespełnieniem warunku czasowego (decyzja h) realizacji projektu pociąga za sobą stratę w wysokości 18 tys. zł wynikającą z tłumaczenia warunków przetargowych oraz częściowego zaangażowania zasobów przedsiębiorstwa w przygotowanie oferty. Stany natury n14 oraz n15 oznaczają natomiast realizację zadania jako dostawca urządzeń, a różnica pomiędzy nimi oznacza procentowy udział produktów przedsiębiorstwa w całości zadania. Przedstawione wyżej dane liczbowe uzyskano na podstawie rzeczywistych założeń projektu oszacowanych dzięki wiedzy i doświadczeniu analityków przedsiębiorstwa odpowiedzialnych za planowanie projektu.

Już na samym początku podejmowania decyzji kierownictwo przedsiębiorstwa zdecydowało ze względów strategicznych na przyjęcie opcji pierwszej, tj. wejścia na nowy zagraniczny rynek kosztem pozyskania małego krajowego kontraktu. W wyniku tej decyzji pominięto analizę liczbową gałęzi dotyczącej realizacji projektu krajowego, a skupiono się na planowaniu działań przedsiębiorstwa na rynku zagranicznym. Wyплаты i prawdopodobieństwa w rozpatrywanym etapie projektu przedstawia tabela 1.

Tabela 1

Wyплаты i prawdopodobieństwa

Etap 1			
Węzeł decyzyjny	Decyzja	Prawdopodobieństwo zdarzenia	
1	a	0.55	
		0.10	
		0.35	
	b	Decyzją kierownictwa pominięto opcję realizacji projektu krajowego	

Etap 2			
Węzeł decyzyjny	Decyzja	Prawdopodobieństwo zdarzenia	Oczekiwana wypłata
2	c	0.75	916
		0.25	-194
	d	0.65	1000
		0.35	-282
	e	0.85	531
		0.15	-104
3	f	1.00	-18
	g	0.60	394
		0.40	-100
4	h	1.00	-18
	i	0.60	462
		0.40	278
5	j	Decyzją kierownictwa pominięto opcję realizacji projektu krajowego	

W kolejnym kroku przeprowadzamy obliczenia stosując metodę maksymalizacji oczekiwanej korzyści. Przykładowy przebieg obliczeń dla węzła decyzyjnego 2 przedstawiono niżej.

Węzeł decyzyjny 2

Decyzje dopuszczalne:

C – oczekiwana wypłata: $0.75 \cdot 916.0 + 0.25 \cdot -194.0 \cong 639$

D – oczekiwana wypłata: $0.65 \cdot 1000.0 + 0.35 \cdot -282.0 \cong 551$

E – oczekiwana wypłata: $0.85 \cdot 531.0 + 0.15 \cdot -104.0 \cong 436$

Decyzja optymalna: C

Zbiorcze wyniki obliczeń przedstawiono w tabeli 2. Kolorem jasnoszarym oznaczono decyzje optymalne w kolejnych węzłach decyzyjnych.

Tabela 2

Wyniki obliczeń oraz decyzje optymalne

Etap 2		
Węzeł decyzyjny	Decyzja	Oczekiwana wypłata
2	C	639
	D	551
	E	436
3	F	-18
	G	196
4	H	-18
	I	388
5	J	Pominięte

Etap 1		
Węzeł decyzyjny	Decyzja	Oczekiwana wypłata
1	A	507
	B	Pominięte

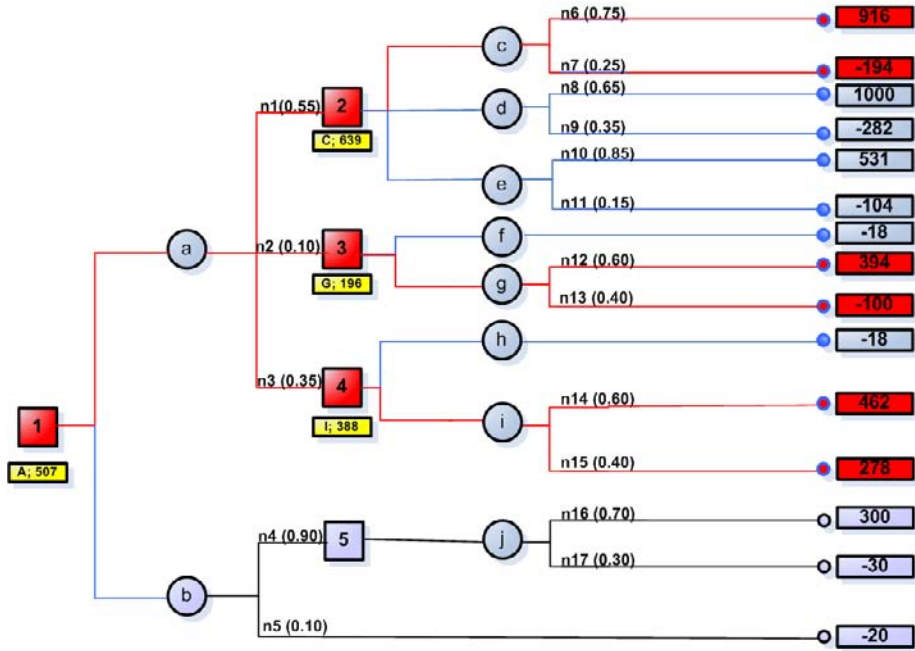
Strategię optymalną przedstawia tabela 3.

Tabela 3

Strategia optymalna

Węzeł decyzyjny	Decyzja
2	C
3	G
4	I
1	A

Omówiony wyżej fragment drzewa decyzyjnego dla obliczeń przeprowadzonych w analizowanym etapie przedstawiono na rys. 2.



Rys. 2. Fragment drzewa decyzyjnego opisującego wybór projektu oraz wybór sposobów jego realizacji po uwzględnieniu wartości poszczególnych ścieżek decyzyjnych

W kolejnym kroku przeprowadzamy eksperymenty symulacyjne. Eksperyment rozpoczynamy od analizy węzła decyzyjnego nr 1. Badanie przeprowadzamy według następującego scenariusza:

1. Na podstawie rozwiązania uzyskanego za pomocą drzewa decyzyjnego określamy decyzję podejmowaną w analizowanym węźle decyzyjnym.

2. Ustalamy węzeł losowy, do którego prowadzi decyzja wyznaczona w punkcie 1.

3. Korzystając z generatora liczb losowych określamy stan natury zachodzący w węźle losowym.

4. Jeżeli zajście określonego stanu natury oznacza, że proces znajdzie się w węźle końcowym, to z rozkładu opisującego korzyść uzyskiwaną w tym węźle losujemy wartość korzyści z realizacji procesu. W przeciwnym wypadku przechodzimy do kolejnego węzła decyzyjnego i wracamy do punktu 1.

Korzystając z modelu symulacyjnego zapisanego w arkuszu kalkulacyjnym EXCEL przeprowadzono 5000 eksperymentów. Uzyskane wyniki pozwalają na konstrukcję rozkładu prawdopodobieństwa zmiennej wyjściowej (marży). Podsumowanie wyników eksperymentów symulacyjnych przedstawiono w tabeli 4.

W wyniku przeprowadzonej symulacji uzyskano następujące wyniki:

- maksymalną stratę w wyniku realizacji projektu: -233,8 tys. zł,
- maksymalny przychód w wyniku realizacji projektu: 955,4 tys. zł.

Tabela 4

Wyniki symulacji (5000 zdarzeń)

Przedział w kPLN	Liczba wylosowanych zdarzeń	Prawdopodobieństwo zajścia zdarzenia
-233,8; 0	793	15,86%
0; 200	141	2,82%
200; 400	1439	28,78%
400; 600	923	18,46%
600; 800	5	0,10%
800; 955,4	1699	33,98%

Podsumowanie

Przeprowadzona symulacja przyjętych w drzewie decyzyjnym strategii wskazała, że:

- w 34% decydent ma możliwość osiągnięcia maksymalnego przychodu z przedsięwzięcia, zatem osiągnięcia maksymalnej marży,
- w 66% decydent nie osiągnie założonego maksymalnego celu finansowego,
- w prawie 16% przedsięwzięcie nie zostanie zrealizowane i przyniesie ewidentną stratę.

Zapoznanie się z wynikami symulacji pozwala decydentowi na wyrobienie szerszej opinii na temat konsekwencji, jakie niesie za sobą przyjęcie strategii postępowania wyznaczonej za pomocą drzewa decyzyjnego.

Zastosowanie drzewa decyzyjnego w połączeniu z symulacją komputerową umożliwiło dokonania analiz nieprzeprowadzanych w środowisku rzeczywistym, gdzie głównym czynnikiem decydującym o przyjęciu lub odrzuceniu określonej strategii jest wyłącznie biznes plan określający możliwy do uzyskania poziom marży. Zastosowanie zatem drzewa decyzyjnego i podjęcie dodatkowego trudu związanego z oszacowaniem prawdopodobieństw zdarzeń losowych wynikających z podejmowanych decyzji umożliwia decydentowi znacznie szybszą reakcję na pojawiające się w trakcie realizacji zdania problemy. Pozwala przewidzieć inne opcje w wyniku niekorzystnego splotu wydarzeń. Pozwala jednak przede wszystkim obracać przy przyjętym prawdopodobnym stanie losowym konkretną strategię decyzyjną.

Przedmiotem dalszych prac będzie próba zaproponowania procedury, która uwzględniałaby możliwość symulacyjnej analizy nie tylko rozwiązania optymalnego z punktu widzenia maksymalizacji oczekiwanej korzyści, ale również rozwiązań suboptymalnych, czyli takich, dla których oczekiwana korzyść jest nieco niższa niż ma to miejsce dla przypadku optymalnego.

Literatura

1. Dąbrowski P. (2003). Kompendium wiedzy o zarządzaniu projektami. Wydawnictwo Management Training & Development Center, Warszawa, 5.
2. Koronacki J., Mielniczuk J. (2007). Statystyka. WNT, Warszawa 2001.
3. Nowak M. (2007). Symulacja komputerowa w problemach decyzyjnych. AE, Katowice.
4. Pritchard C. (2001). Zarządzanie ryzykiem w projektach. WIG-Press, Warszawa.
5. Schuyler J. (2001). Risk and Decision Analysis In Project. PMI, Atlanta.
6. Trzaskalik T. (2003). Wprowadzenie do badań operacyjnych z komputerem. PWE, Warszawa.

Sebastian Sitarz

Uniwersytet Śląski w Katowicach

OBJĘTOŚCIOWA ANALIZA WRAŻLIWOŚCI W PROGRAMOWANIU LINIOWYM

Wprowadzenie

Badania nad analizą wrażliwości oraz programowaniem parametrycznym w zadaniach optymalizacji liniowe datują się od pracy Manne'a [7], który zajmował się zmianami wektora wyrazów wolnych. Natomiast terminu „analiza wrażliwości” jako pierwszy użył Heller [6]. Dalszy rozwój badań w dziedzinie wrażliwości rozgałęził się na wiele różnych zagadnień. Warto wspomnieć o monografii Banka i innych [1] opisującą analizę wrażliwości w zagadnieniach nieliniowej optymalizacji w przestrzeniach metrycznych. Innym interesującym przykładem rozwoju analizy wrażliwości jest praca Wendella [11], w której wprowadza się tzw. maksymalną tolerancję zmian parametrów w zadaniach programowania liniowego (zmiany dotyczą współczynników funkcji celu, ograniczeń i wektora wyrazów wolnych).

Warto nadmienić również o pracach dotyczących analizy wrażliwości w zadaniach wielokryterialnego programowania liniowego. Jedną z pierwszych prac poświęconych tej tematyce jest artykuł Hansena i innych [5], w którym autorzy przeprowadzają analizę wrażliwości po skalaryzacji wyjściowego zadania wielokryterialnego. Analizą wrażliwości w zadaniach programowania celowego przeprowadzili Dauer i Liu [2]. W pracy Sitarza [10] znajdujemy metody wyznaczania przedziałów współczynników funkcji celu w wielokryterialnym programowaniu liniowym.

Szerszy przegląd współczesnych prac nad problemami postoptymalizacyjnymi znajdziemy w pracy Gala i Greenberga [4].

Niniejsza praca opisuje podejście objętościowe zaproponowane przez Vetschera [12], które zostawia jednak pewne luki dotyczące aspektów obliczeniowych metody. W związku z powyższym celem pracy jest przedstawienie pełnego opisu metody uwzględniającego aspekt numeryczny (w szczególności sposób wyznaczania objętości wielościanu). Wykorzystane tu są wyniki otrzymane w pracy magisterskiej Joanny Mazurek [8].

1. Podstawowe pojęcia i notacja

1.1 Zadanie programowania liniowego

W dalszym ciągu będziemy rozważać zadanie programowania liniowego w postaci

$$\text{Max } \{ \mathbf{c}\mathbf{x} : \mathbf{x} \in X \} \quad (1)$$

gdzie:

$X = \{ \mathbf{x} \in \mathbb{R}^n : \mathbf{A}\mathbf{x} \leq \mathbf{b}, \mathbf{x} \geq 0 \}$ lub $X = \{ \mathbf{x} \in \mathbb{R}^n : \mathbf{A}\mathbf{x} = \mathbf{b}, \mathbf{x} \geq 0 \}$ – zbiór rozwiązań dopuszczalnych,

$\mathbf{c} \in \mathbb{R}^n$ – wektor współczynników funkcji celu,

$\mathbf{A} \in \mathbb{R}^{n,m}$ – macierz współczynników ograniczeń rzędu m ,

$\mathbf{b} \in \mathbb{R}^m$ – wektor ograniczeń,

$\mathbf{x} \in \mathbb{R}^n$ – wektor zmiennych.

1.2. Zapis tablicy simpleks

Dla zadania (1) ze zbiorem rozwiązań dopuszczalnych

$$X = \{ \mathbf{x} \in \mathbb{R}^n : \mathbf{A}\mathbf{x} = \mathbf{b}, \mathbf{x} \geq 0 \}$$

przyjmujemy następujące oznaczenia:

$\mathbf{A}_B = [\mathbf{a}^{j_1}, \mathbf{a}^{j_2}, \dots, \mathbf{a}^{j_m}]$ – podmacierz macierzy $\mathbf{A}$, taka, że $\mathbf{a}^{j_1}, \mathbf{a}^{j_2}, \dots, \mathbf{a}^{j_m}$ tworzą bazę $\mathbb{R}^m$

$\mathbf{B} = \{ j_1, \dots, j_m \}$ – zbiór indeksów bazowych

$\mathbf{A}_N$ – podmacierz macierzy $\mathbf{A}$ bez kolumn bazowych

$\mathbf{x}$ – dopuszczalne rozwiązanie bazowe związane z $\mathbf{B}$, gdy $\mathbf{x}_B = \mathbf{A}_B^{-1} \mathbf{b} \geq 0$ oraz $\mathbf{x}_N = 0$.

$\mathbf{d} = \mathbf{c} - \mathbf{c}_B \mathbf{A}_B^{-1} \mathbf{A}$ – wskaźniki optymalności.

Szablon stosowanej tu tablicy simpleks związanej a bazą $\mathbf{B}$ przedstawia tabela 1.

Tabela 1

Tablica simpleks związana z bazą $\mathbf{B}$

$\mathbf{x}$	
$\mathbf{c}$	
$\mathbf{A}_B^{-1} \mathbf{A}$	$\mathbf{A}_B^{-1} \mathbf{b}$
$\mathbf{d}$	

1.3. Analiza wrażliwości

Analiza wrażliwości danego rozwiązania optymalnego ze względu na zmiany współczynników funkcji celu polega na zbadaniu dopuszczalnych zmian współczynników w taki sposób, aby dane rozwiązanie pozostawało rozwiązaniem optymalnym. Zgodnie z tą ideą rozważamy zmodyfikowane zadanie (1) postaci

$$\text{Max } \{c(t) \mid x \in X\} \quad (2)$$

gdzie:

t – wektor parametrów.

Wektor wskaźników optymalności zadania (2) oznaczać będziemy przez $d(t)$.

Uwaga 1. Warto dodać, że w literaturze dotyczącej analizy wrażliwości, współczynniki wektora funkcji celu są zapisywane często bezpośrednio jako parametry (bez użycia t). Ta uwaga ułatwia zapis i będzie również stosowana w niniejszej pracy.

1.4. Standardowa analiza wrażliwości

Standardowa analiza wrażliwości dotyczy zmian wyłącznie jednego współczynnika funkcji celu (bez zmian pozostałych). Zgodnie z powyższym, wektor $c(t)$ jest postaci

$$c(t) = c + t$$

gdzie $t = [0, 0, \dots, t_i, 0, \dots, 0]$.

Przypadek ten został dokładnie opisany w monografii Gala [3].

2. Objętościowa analiza wrażliwości

Standardowe podejście analizy wrażliwości jest szeroko znane, jednak jest ograniczone co do zmiany wyłącznie jednego współczynnika (ma to związek z łatwą interpretacją otrzymanego wyniku – przedziału zmienności współczynnika). Stosowanie analizy wrażliwości przy zmianach wszystkich współczynników funkcji celu jest ogólniejsze, jednak interpretacja otrzymanego wyniku jest trudna (z reguły otrzymuje jako wynik obszerny zbiór współczynników, trudny w interpretacji). Celem objętościowej analizy wrażliwości jest

połączenie powyższych podejść. Będziemy jednocześnie zmieniać wszystkie współczynniki, a jako końcowy wynik otrzymamy jedną liczbę będącą miarą wrażliwości – objętość zbioru dopuszczalnych zmian współczynników (po unormowaniu współczynników). Warto dodać, że podobne podejście zaproponował Wendel [11]. Różnica polega na tym, że w podejściu Wendella znajdujemy maksymalny promień kuli leżącej wewnątrz zbioru dopuszczalnych współczynników.

2.1. Wyznaczanie zbioru dopuszczalnych zmian wszystkich współczynników

Wyznaczanie zmian współczynników, dla których dane rozwiązanie jest optymalne, sprowadza się do rozwiązania układu nierówności liniowych (wskaźniki optymalności powinny być niedodatnie)

$$\mathbf{d}(\mathbf{t}) \leq \mathbf{0}$$

W pracy Vetschera [12] znajdujemy alternatywny sposób wyznaczania zmian współczynników – metodę sąsiednich baz.

Metoda sąsiednich baz

W celu sformułowania metody wprowadzimy niezbędne oznaczenia i definicje:

$k(1), k(2), \dots, k(n-m)$ – indeksy niebazowych zmiennych,

$\mathbf{B}^{k(j)}$ – baza sąsiednia do bazy $\mathbf{B}$ powstała przez wprowadzenie zmiennej o indeksie $k(j)$,

x_l^{opt} – wartość zmiennej x_l w wyjściowej, optymalnej bazie $\mathbf{B}$,

$x_{k(j),l}$ – wartość zmiennej x_l w bazie $\mathbf{B}^{k(j)}$,

$\overset{c}{\mathbf{X}}$ – macierz kwadratowa stopnia $n-m$ złożona z następujących elementów

$$c_{l,k(j)} = x_l^{opt} - x_{k(j),l} \quad \text{dla } l, k = 1, 2, \dots, n-m$$

Tablicą baz sąsiednich nazywamy tablicę przedstawioną w tabeli 2. Zwróćmy uwagę, że w tablicy baz sąsiednich kolumny pod współczynnikami c_i tworzą macierz $\overset{c}{\mathbf{X}}$. Zgodnie z powyższymi definicjami, tablica baz sąsiednich służy do prezentacji wartości zmiennych w sąsiednich bazach oraz prezentacji macierzy $\overset{c}{\mathbf{X}}$.

Tabela 2

Tablica baz sąsiednich

Indeksy niebazowe	Nowe wartości			Współczynniki		
	x_1	...	x_{n-m}	c_1	...	c_{n-m}
$k(1)$	$x_{k(1),1}$	...	$x_{k(1),n-m}$	$c_{1,k(1)}$	...	$c_{n-m,k(1)}$
...	...	...	...	...	...	...
$k(n-m)$	$x_{k(n-m),1}$	...	$x_{k(n-m),n-m}$	$c_{1,k(n-m)}$	...	$c_{n-m,k(n-m)}$

W pracy Vetchera [12] podano twierdzenie, wedle którego współczynniki wektora c , dla których dane rozwiązanie jest optymalne spełniają zależność

$$\bar{\mathbf{X}}^c \mathbf{c} \geq 0$$

Powyższa zależność określa zbiór, który w przypadku parametryzacji zgodnej z uwagą 1 jest wielościanem wypukłym generowanym przez nierówności liniowe. Takim przypadkiem parametryzacji zajmujemy się dalej.

2.2. Obliczanie objętości

Objętościowa analiza wrażliwości opiera się na zmierzeniu objętości wielościanu określonego przez ograniczenia $\bar{\mathbf{X}}^c \mathbf{c} \geq 0$. W ogólnym przypadku objętość ta jest nieskończona, więc autor metody [12] proponuje znormalizowanie nieujemnych współczynników.

Zatem w dalszym ciągu będziemy zajmować się liczeniem objętości zbioru opisanego następująco

$$W = \left\{ \mathbf{c} \in R^n : \bar{\mathbf{X}}^c \mathbf{c} \geq 0 \wedge \mathbf{c} \geq 0 \wedge \sum_{i=1}^n c_i = 1 \right\}$$

Normalizacja nieujemnych współczynników funkcji celu w programowaniu liniowym ma na celu:

- ograniczenie się do badania wrażliwości w simpleksie jednostkowym,
- możliwość porównywania wyników między różnymi zadaniami PL.

W przypadku nieujemnych współczynników normalizacja nie ogranicza stosowalności. Istotnie, jeśli współczynniki nie sumują się do jedynki, to ich znormalizowanie (pomnożenie przez dodatni skalar) przekształca wyjściowe zadanie w równoważne zadanie optymalizacyjne. Natomiast w metodzie Vetschery nie dopuszcza się stosowania ujemnych współczynników.

Przedstawiona tu metoda obliczania objętości powyższego wielościanu jest oparta na metodzie Monte Carlo (MMC), czyli konstrukcji pomocniczego procesu losowego. Będziemy używać MMC do obliczenia objętości wielościanu W zawierającego się w jednostkowym simpleksie. Proponowana metoda polega na losowaniu punktu z simpleksu jednostkowego (szczegółowy opis znajdziemy w [9]). Następnie należy sprawdzić, czy wylosowany punkt spełnia zależności: $\sum_{i=1}^n c_i \geq 0$. Procedurę losowania i sprawdzania powtarzamy wielokrotnie zgodnie z metodą MMC. Przyjmując oznaczenia:

N – liczba wylosowanych punktów z simpleksu jednostkowego,

L – liczba wylosowanych punktów spełniających warunek $\sum_{i=1}^n c_i \geq 0$,

objętość wielościanu W jest równa w przybliżeniu $\frac{L}{N} \cdot \frac{1}{(n-1)!}$, gdyż $\frac{1}{(n-1)!}$

jest objętością simpleksu: $\left\{ c \in R^n : c \geq 0 \wedge \sum_{i=1}^n c_i = 1 \right\}$.

3. Przykład

Rozważmy następujące zadanie programowania liniowego pochodzące z pracy [12]:

$$\begin{aligned} & \text{Max } 0,029x_1 + 0,35x_2 + 0,27x_3 + 0,35x_4 \\ & \quad 0,95x_1 + 0,23x_2 + 0,79x_3 + 0,40x_4 \leq 0,67 \\ & \quad 0,51x_1 + 0,42x_2 + 0,62x_3 + 0,27x_4 \leq 0,29 \\ & \quad 0,66x_1 + 0,08x_2 + 0,31x_3 + 0,10x_4 \leq 0,95 \\ & \quad 0,32x_1 + 0,11x_2 + 0,52x_3 + 0,65x_4 \leq 0,10 \\ & \quad 0,87x_1 + 0,55x_2 + 0,54x_3 + 0,29x_4 \leq 0,39 \\ & \quad x_1, x_2, x_3, x_4 \geq 0 \end{aligned}$$

Przeprowadzimy teraz analizę wrażliwości dotyczącą tego zadania, przekraczając wyniki uzyskane w pracy Vetschera [12]. Po sprowadzeniu powyższego zadania do postaci równościowej otrzymujemy rozwiązanie optymalne przedstawione w tabeli 3 odpowiadające bazie $B=\{2, 4, 5, 7, 9\}$.

Tabela 3

Tablica simpleks rozwiązania optymalnego

x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	
0,029	0,350	0,270	0,350	0	0	0	0	0	
0,590	0	0,295	0	1	-0,434	0	-0,435	0	0,501
1,007	1	1,079	0	0	2,672	0	0	0	0,664
0,547	0	0,162	0	0	-0,169	1	0	0	0,893
0,322	0	0,617	1	0	-0,452	0	1	0	0,042
0,223	0	-0,233	0	0	-1,338	0	0	1	0,012
-0,436	0	-0,324	0	0	-0,777	0	-0,216	0	

Sparametryzujemy wyjściowe współczynniki funkcji celu (0, 029; 0, 35; 0, 27; 0, 35) wektorem (c_1, c_2, c_3, c_4) i zastosujemy metodę bazy sąsiedniej w celu znalezienia warunków określających zmiany współczynników funkcji celu, przy których otrzymane rozwiązanie pozostaje optymalnym rozwiązaniem.

Zastosujemy teraz metodę baz sąsiednich. Wprowadzając do optymalnej bazy kolejne zmienne niebazowe: x_1, x_3, x_6, x_8 , otrzymamy odpowiednio tabele 4, 5, 6, 7.

Tabela 4

Tablica simpleks odpowiadająca bazie B^1

x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	
0,029	0,350	0,270	0,350	0	0	0	0	0	
0	0	0,911	0	1	3,106	0	-0,726	-2,646	0,469
0	1	2,131	0	0	8,714	0	-1,608	-4,516	0,610
0	0	0,734	0	0	3,113	1	-0,354	-2,453	0,864
0	0	0,953	1	0	1,480	0	1,567	-1,444	0,0247
1	0	-1,045	0	0	-6,010	0	0,493	4,484	0,054
0	0	-0,779	0	0	-3,394	0	-0,0001	1,956	

Tabela 5

Tablica simpleks odpowiadająca bazie B^3

x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	
0,029	0,350	0,270	0,350	0	0	0	0	0	
0,436	0	0	-0,478	1	-0,218	0	-1,260	0	0,481
0,444	1	0	-1,749	0	3,462	0	-4,129	0	0,591
0,462	0	0	-0,263	0	-0,050	1	-0,537	0	0,882
0,522	0	1	1,621	0	-0,733	0	2,797	0	0,068
0,345	0	0	0,378	0	-1,509	0	0,762	1	0,028
-0,267	0	0	0,524	0	-1,014	0	0,690	0	

Tabela 6

Tablica simpleks odpowiadająca bazie $\mathbf{B}^6$

x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	
0,029	0,350	0,270	0,350	0	0	0	0	0	
0,754	0,162	0,470	0	1	0	0	-0,615	0	0,609
0,377	0,374	0,404	0	0	1	0	-0,416	0	0,249
0,611	0,063	0,230	0	0	0	1	-0,154	0	0,935
0,492	0,169	0,800	1	0	0	0	1,538	0	0,154
0,727	0,501	0,307	0	0	0	0	-0,446	1	0,344
-0,143	0,291	-0,010	0	0	0	0	-0,538	0	

Tabela 7

Tablica simpleks odpowiadająca bazie $\mathbf{B}^8$

x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	
0,029	0,350	0,270	0,350	0	0	0	0	0	
0,671	0	0,451	0,252	1	-0,548	0	0	0	0,512
1,214	1	1,476	0,644	0	2,381	0	0	0	0,691
0,563	0	0,192	0,049	0	-0,191	1	0	0	0,895
0,187	0	0,357	0,579	0	-0,262	0	1	0	0,024
0,202	0	-0,272	-0,064	0	-1,309	0	0	1	0,009
-0,396	0	-0,247	0,125	0	-0,833	0	0	0	

Wykorzystując tabele 4-7 otrzymujemy dla bazy $\mathbf{B} = \{2, 4, 5, 7, 9\}$ tablicę baz sąsiednich przedstawioną w tabeli 8.

Tabela 8

Tablica baz sąsiednich dla bazy $\mathbf{B} = \{2, 4, 5, 7, 9\}$

Indeksy niebazowe	Nowe wartości				Współczynniki			
	x_1	x_2	x_3	x_4	c_1	c_2	c_3	c_4
1	0,054	0,610	0	0,025	-0,054	0,054	0	0,017
3	0	0,591	0,068	0	0	0,073	-0,068	0,042
6	0	0,000	0	0,154	0	0,664	0	-0,112
8	0	0,691	0	0	0	-0,027	0	0,042

Przykładowo, pokażemy jak obliczono w tablicy baz sąsiednich wartości: $x_{8,2} = 0,691$ oraz $c_{8,2} = -0,027$. Otóż wartość $x_{8,2}$ jest to wartość zmiennej o indeksie 2 w bazie $\mathbf{B}^8$ (tabela 7). Natomiast wartość $c_{8,2}$ jest równa różnicy wartości zmiennej x_2 w optymalnym rozwiązaniu (tabela 3) i wartości $x_{8,2}$

$$c_{l,k(j)} = x_l^{opt} - x_{k(j),l} = x_2^{opt} - x_{8,2} = 0,664 - 0,691 = -0,027.$$

Macierz $\overset{c}{\mathbf{X}}$ odczytana z tablicy baz sąsiednich wynosi

$\overset{c}{\mathbf{X}} =$	-0,054	0,054	0	0,017
	0	0,073	-0,068	0,042
	0	0,664	0	-0,112
	0	-0,027	0	0,042

Zatem zależności $\overset{c}{\mathbf{X}} \mathbf{c} \geq 0$ mają postać

$$\begin{aligned} -0,054c_1 + 0,054c_2 + \quad \quad \quad + 0,017c_4 &\geq 0 \\ 0,073c_2 - 0,068c_3 + 0,042c_4 &\geq 0 \\ 0,664c_2 \quad \quad \quad - 0,112c_4 &\geq 0 \\ -0,027c_2 \quad \quad \quad + 0,042c_4 &\geq 0 \end{aligned}$$

Przejdziemy teraz do obliczenia objętości wielościanu

$$W = \left\{ \mathbf{c} \in \mathbb{R}^n : \overset{c}{\mathbf{X}} \mathbf{c} \geq 0 \wedge \mathbf{c} \geq 0 \wedge \sum_{i=1}^4 c_i = 1 \right\}$$

Zgodnie z metodą opisaną w rozdziale 2 przyjęto $N=10\,000$. Liczba wylosowanych punktów z simpleksu jednostkowego spełniających warunków $\overset{c}{\mathbf{X}} \mathbf{c} \geq 0$ wyniosła $L=2\,149$. Zatem objętość rozpatrywanego wielościanu wynosi

$$\frac{L}{N} \cdot \frac{1}{(n-1)!} = \frac{2\,149}{10\,000} \cdot \frac{1}{3!} = 0,0358.$$

Na otrzymany wynik możemy też spojrzeć następująco: otóż prawdopodobieństwo, że rozpatrywane rozwiązanie będzie optymalne (dla losowo wybranych współczynników z simpleksu jednostkowego) wynosi: $\text{prob}=0,21489$.

Uwaga 2. Wykorzystamy teraz prezentowane podejście objętościowe w następującym problemie: dla każdej danej dopuszczalnej bazy wyznaczyć prawdopodobieństwo (prob), że dana baza będzie optymalna.

W tym celu należy przeprowadzić procedurę wyznaczania zależności $\overset{c}{\mathbf{X}} \mathbf{c} \geq 0$ dla każdej z baz, a następnie szacować objętości jak opisano w rozdziale 2.2. Przeprowadzając powyższe procedury, dla wszystkich dopuszczalnych baz, otrzymujemy prawdopodobieństwa przedstawione w tabeli 9.

Tabela 9

Prawdopodobieństwa optymalności dla wszystkich dopuszczalnych baz

baza	prob
{ 1, 5, 6, 7, 9 }	0,1896
{ 2, 5, 7, 8, 9 }	0,2412
{ 3, 5, 6, 7, 9 }	0,0928
{ 4, 5, 6, 7, 9 }	0,0598
{ 1, 2, 5, 7, 9 }	0,1077
{ 2, 3, 5, 7, 9 }	0,0941
{ 2, 4, 5, 7, 9 }	0,2149

Spójrzmy na tabelę 9 z punktu widzenia programowania parametrycznego, w którym bada się zmienność rozwiązań optymalnych ze względu na zmienność parametrów. W klasycznym podejściu programowania parametrycznego przeszukujemy wartości parametrów, podając odpowiadające im rozwiązania optymalne. Natomiast w podejściu objętościowym podajemy objętości zbioru parametrów, dla których dane rozwiązanie jest optymalne.

Podsumowanie i dalsze kierunki badań

W pracy pokazano podejście objętościowe w analizowaniu wrażliwości współczynników funkcji celu w zadaniach programowania liniowego. Ponadto, przedstawiono schemat numeryczny służący do wyznaczenia objętości zbioru składającego się ze współczynników funkcji celu, dla których dane rozwiązanie pozostaje optymalne. Warto zauważyć, że metoda wykorzystująca tablicę baz sąsiednich nie daje korzyści numerycznych w porównaniu z wykorzystaniem układu nierówności $\mathbf{d}(\mathbf{t}) \leq \mathbf{0}$. Natomiast samo podejście Vetschery pokazuje ciekawe teoretyczne związki w programowaniu liniowym.

Dalsze kierunki badań nad omawianą tematyką, stanowią następujące zagadnienia:

- porównanie podejścia objętościowego z tolerancyjną analizą wrażliwości zaproponowaną przez Wendella,
- zastosowanie podejścia objętościowego w zadaniach wielokryterialnego programowania liniowego (wyznaczanie objętości wielościanu współczynników, dla których dane rozwiązanie jest rozwiązaniem sprawnym),
- stworzenie algorytmu wyznaczającego dokładną objętość otrzymywanych wielościanów (bez użycia metod losowych).

Literatura

1. Bank B., Guddat D., Klatte B., Kummer K. (1982). Nonlinear Parametric Optimization. Akademie Verlag, Berlin.
2. Dauer J., Liu Y.H. (1997). Multiple-Criteria and Goal Programming. In: Advances in Sensitivity Analysis and Parametric Programming. Section 11. Eds. T. Gal, H.J. Greenberg. Kluwer Academic Publishers, Boston.
3. Gal T. (1979). Postoptimal Analyses, Parametric Programming and Related Topics. Walter de Gruyter, Berlin.
4. Gal T., Greenberg H.J. (1997). Advances in Sensitivity Analysis and Parametric Programming. Kluwer Academic Publishers, Boston.
5. Hansen P., Labbe M., Wendell R.E. (1989). Sensitivity Analysis in Multiple Objective Linear Programming: The Tolerance Approach. European Journal of Operational Research, 38, 63-69.
6. Heller I. (1954). Sensitivity Analysis in Linear Programming. Logistics Research Project, The George Washington University.
7. Manne A.S. (1953). Notes on Parametric Linear Programming. RAND-Corporation Report, No p-468.
8. Mazurek J. (2008). Objętościowa analiza wrażliwości. Uniwersytet Śląski, Katowice.
9. Robert C., Casella G. (1999). Monte Carlo Statistical Methods. Springer-Verlag, Berlin.
10. Sitarz S. (2008). Postoptimal Analysis in Multicriteria Linear Programming. European Journal of Operational Research, 191/1, 7-18.
11. Vetschera R. (1997). Volume-Based Sensitivity Analysis for Multi-Criteria Decision Models. In: Methods of Multicriteria Decision Theory. Eds. A. Göpfert, J. Seeländer, Chr. Tammer. Hänsel-Hohenhausen, Egelsbach.
12. Wendell R.E. (1982). A Preview of a Tolerance Approach to Sensitivity Analysis in Linear Programming. Discrete Mathematics, 38, 121-124.

Adam Sojda

Politechnika Śląska w Gliwicach

WYKORZYSTANIE ALGORYTMU EWOLUCYJNEGO W ZAGADNIENIU LOKALIZACJI SYSTEMÓW MASOWEJ OBSŁUGI G/G/1

Wprowadzenie

W pracy przedstawiono zastosowanie algorytmu ewolucyjnego do zagadnienia znalezienia lokalizacji p obiektów będących systemami masowej obsługi G/G/1. Systemy te służą do obsługi n obiektów. Znanych jest m potencjalnych lokalizacji systemów masowej obsługi, jak i również znana jest lokalizacja obsługiwanych obiektów wraz z siecią połączeń.

Zagadnienie lokalizacji jest zagadnieniem związanym z ustaleniem przynależności obiektów zwanych klientami do obiektów zwanych centrami dystrybucji. Przeważnie staramy się rozmieścić systemy obsługi tak, aby znajdowały się jak najbliżej obiektów obsługiwanych. Rozmieszczenie to powinno uwzględniać znaną sieć połączeń drogowych, telekomunikacyjnych itp. Dla pewnej klasy obiektów zwanej centrami serwisowymi (policja, pogotowie ratunkowe, straż, pogotowie gazowe, energetyczne) centra takie należy zlokalizować w taki sposób, aby czas dostępu i obsługi przez centrum był jak najmniejszy. W takich przypadkach centra serwisowe są modelowane za pomocą systemów masowej obsługi, a metody wyznaczenia rozwiązania można znaleźć w pracach [1; 2; 4].

1. Założenia modelu

W systemach masowej obsługi oznaczonych jako G/G/1 zgłoszenia przybywają do stanowiska obsługi w niezależnych odstępach czasu z wartością oczekiwaną $\frac{1}{\lambda}$ oraz wariancją σ_A^2 . Czas obsługi zgłoszenia w systemie jest

zmienną losową o wartości oczekiwanej $\frac{1}{\mu}$ oraz wariancji σ_B^2 . Dla takiego systemu na podstawie równań procesu dyfuzji można otrzymać funkcję gęstości a następnie wyznaczyć parametry systemu [3; 5; 6].

Dane są zbiory:

$I = \{1, 2, \dots, n\}$ – numerów obiektów w sieci,

$J = \{1, 2, \dots, m\}$ – numerów potencjalnych lokalizacji systemów.

Dane są również następujące parametry:

t_{ij} – średni czas przejazdu pomiędzy i -tym obiektem a j -tą lokalizacją
 $i \in I, j \in J$,

$\frac{1}{\mu_j}$ – średni czas obsługi w j -tym systemie,

$\sigma_{B,j}^2$ – wariancja czasu obsługi w j -tym systemie,

λ_i – średnia liczba zgłoszeń w jednostce czasu dla i -tego obiektu,

$\sigma_{A,i}^2$ – wariancja rozkładu odstępów czasu pomiędzy kolejnymi zgłoszeniami z obiektu i -tego,

p – liczba zakładanych systemów.

Można wyznaczyć parametry dla całego systemu przy ustalonej lokalizacji i przyporządkowaniu obiektów do lokalizacji.

Parametry:

średnia liczba zgłoszeń w jednostce czasu do j -tego systemu

$$\lambda^{(j)} = \sum_{i \in S_j} \lambda_i \quad (1)$$

parametr

$$\rho_j = \frac{\lambda^{(j)}}{\mu_j} < 1 \quad (2)$$

oraz

$$c_A^{2(j)} = \frac{1}{\lambda^{(j)}} \sum_{i \in S_j} \lambda_i^3 \sigma_{Ai}^2 \quad (3)$$

$$c_B^{2(j)} = \sigma_{Bj}^2 \mu_j^2 \quad (4)$$

Średni czas obsługi zgłoszenia w j -tym systemie wyraża się wzorem

$$T_j = \frac{\left(0,5 + \frac{c_A^{2(j)} \rho_j + c_B^{2(j)}}{2(1 - \rho_j)} \right) \rho_j}{\lambda^{(j)}} \quad (5)$$

Zmienne decyzyjne

$$x_{ij} = \begin{cases} 1 & \text{gdy } i\text{-ty obiekt jest przypisany do } j\text{-tego systemu} \\ 0 & \text{w przeciwnym przypadku} \end{cases} \quad (6)$$

$$y_j = \begin{cases} 1 & \text{gdy system zlokalizowano w } j\text{-tej lokalizacji} \\ 0 & \text{w przeciwnym przypadku} \end{cases} \quad (7)$$

Ograniczenia

lokalizacja p systemów

$$\sum_{j \in J} y_j \leq p \quad (8)$$

każdy obiekt jest obsługiwany przez jeden system

$$\sum_{j \in J} x_{ij} = 1, \quad i \in I \quad (9)$$

każdy obiekt obsługiwany jest przez założony system obsługujący

$$x_{ij} \leq y_j, \quad i \in I, \quad j \in J \quad (10)$$

wiadomo również, że średni czas obsługi w systemie jest równy

$$T_{ij} = (T_j + t_{ij})x_{ij}, \quad i \in I, \quad j \in J \quad (11)$$

gdzie wartość T_{ij} wyznaczany jest na podstawie wzoru (5) z uwzględnieniem wzorów (3) i (4).

W zagadnieniu lokalizacji za kryterium przyjmuje się następujące funkcje K1: minimalizację średniego czasu realizacji wszystkich klientów

$$\sum_{i \in I} \sum_{j \in J} T_{ij} \rightarrow \min \quad (12)$$

K2: minimalizację maksymalnego czasu obsługi systemu obsługi

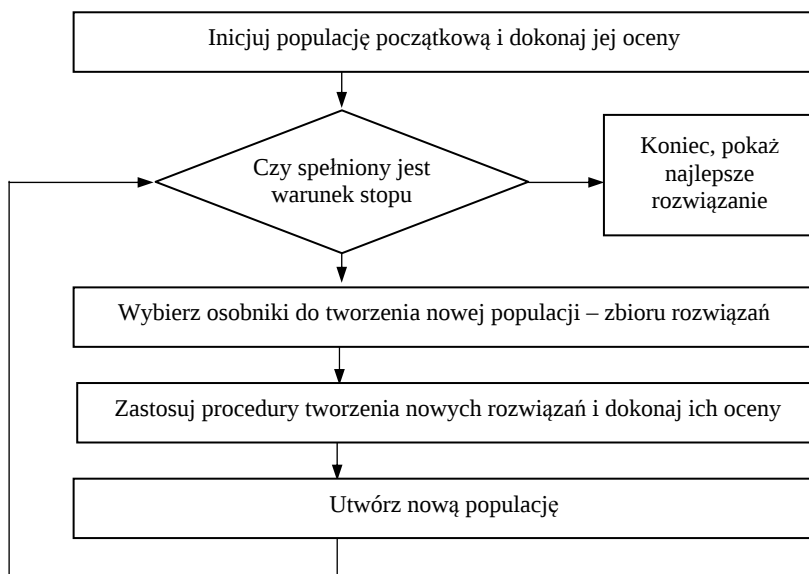
$$\max_{i \in S_j} T_{ij} \rightarrow \min \quad (13)$$

2. Algorytm ewolucyjny

Algorytm ewolucyjny może być utożsamiany z pewnym środowiskiem, w którym określone są mechanizmy generowania nowych rozwiązań. Do klasycznych mechanizmów zalicza się krzyżowanie oraz mutację. Warto zaznaczyć, że analizowany jest pewien zbiór rozwiązań zwany populacją, zatem w środowisku zaimplementowane są mechanizmy przeszukiwania równoległego przestrzeni rozwiązań. Rozwój metod ewolucyjnych był związany z rozwojem możliwości obliczeniowych, co sprawiło, że możemy do rozwiązywania problemów używać analogii biologicznych budując środowiska,

w których rozwiązania mogą, zgodnie ze zdefiniowanymi regułami, ewoluować. Za najpopularniejszy algorytm ewolucyjny należy niewątpliwie uznać algorytm genetyczny – wprowadzony przez Hollanda w 1975 roku i modyfikowany przez innych badaczy [7; 8; 11]. Do mniej znanych algorytmów ewolucyjnych zalicza się ewolucja różnicowa [9; 10].

Ogólny schemat algorytmu ewolucyjnego przedstawia rys. 1.



Rys. 1. Schemat algorytmu ewolucyjnego

Metody ewolucyjne posługują się niedeterministyczną procedurą tworzenia nowych rozwiązań, co nie zawsze może prowadzić do wygenerowania rozwiązania dopuszczalnego jak i optymalnego.

Dla rozważanego zagadnienia rozwiązanie zakodowane jest za pomocą skończonego ciągu, wektora o długości $n + p$. Na pierwszych n miejscach mogą pojawiać się wartości $1, 2, \dots, p$, gdyż tyle mamy możliwości lokalizacji. Następne p elementów określa poszczególne centra, ponieważ wybieramy z m lokalizacji, zatem wartości pojawiające się na pozycjach $n + 1$ do $n + p$ to $1, 2, \dots, m$. Interpretacja symboli jest dynamiczna. Na początku dokujemy interpretacji lokalizacji centrów. Wartości pojawiające się na pozycjach $n + 1$ do $n + p$ powinny być różne. Jeśli tak nie jest, kolejnym powtórzeniom tego samego symbolu nadajemy inne znaczenie. Jeśli dany symbol się powtórzył, to zamieniamy go na następny w kolejności, który nie został jeszcze użyty. Taka procedura pozwala na jednoznaczne wyznaczenie dokładnie p potencjalnych lokalizacji, choć kosztem nadmiarowości. Następnie należy zinterpretować symbole pojawiające się na pierwszych n miejscach. Wartości te interpretujemy

jako wskaźniki, z której pozycji podciągu odpowiadającego za centra należy odczytać numer lokalizacji. Przykładowo, jeśli na trzecim miejscu znajduje się wartość 2, to w podciągu odpowiedzialnym za lokalizację centrów odczytujemy jaką lokalizacja jest przypisana do pozycji $n + 2$. Zatem obiekt trzeci będzie obsługiwany przez centrum określone pozycją $n + 2$.

W stosunku do zaproponowanej pierwotnej wersji ewolucji różnicowej [9; 10] należy wprowadzić zasadnicze zmiany. Rozwiązanie – wektor reprezentowane będzie za pomocą skończonego ciągu zawierającego liczby naturalne od 1 do p oraz od 1 do m zastępując reprezentację zmiennoprzecinkową klasycznego algorytmu ewolucji różnicowej. Wszystkie procedury dla algorytmu ewolucji różnicowej są zachowane z uwzględnieniem powyższej modyfikacji. Tworzenie wektora różnicowego i jego skalowanie odbywa się na zasadzie wykonywania działań w grupie Z_m bądź Z_p z przesunięciem o 1. Dla każdego z rozwiązań bieżącej populacji x_i wybierane są losowo trzy rozwiązania – wektory x_{Ai} , x_{Bi} , x_{Ci} . Następnie tworzony jest wektor x_{Di} zgodnie ze wzorem

$$x_{Di} = (x_{Ci} + [F \cdot (x_{Ai} + x_{Bi})])_{m \text{ albo } p} + 1 \quad (14)$$

gdzie F oznacza stałą skalowania, $[x]$ – całość z x , $(x)_{m \text{ albo } p}$ – resztę z dzielenia x przez m albo p , czyli $x \bmod m$ albo $x \bmod p$. Otrzymany w ten sposób wektor jest następnie krzyżowany jednorodnie z wektorem x_i , celem otrzymania wektora $x'_i = x_i \otimes x_{Di}$. Krzyżowanie jednorodne polega na wygenerowaniu wartości losowej z przedziału od 0 do 1. Jeśli wygenerowana wartość jest mniejsza od założonej wartości prawdopodobieństwa p_c , nowe rozwiązanie tworzone jest na podstawie pierwszego składnika, jeśli większa, udział w tworzeniu rozwiązania ma składnik drugi. Powstać mogą dwa rozwiązania, których składowe pochodzą raz od jednego raz od drugiego „rodzica”. Rozwiązanie x_i jest porównywane z rozwiązaniem x'_i oraz x''_i . Lepsze z nich zajmuje i -tą pozycję w następnej populacji.

Dodatkowo zaproponowano, aby dla danej lokalizacji poszukać najlepszego przydziału na podstawie istniejących już rozwiązań. Dla i -tego rozwiązania uruchamiany jest podrzędny algorytm ewolucji różnicowej bazujący na tej samej populacji, jednakże ograniczonej do podciągu odpowiadającego za przydział klientów do centrum. Do następnego pokolenia przechodzi najlepsze ze znalezionych rozwiązań.

Ze względu na dyskretny charakter zadania proponuje się wprowadzenie heurystycznej procedury poszukiwawczej, tworząc hybrydowy algorytm ewolucyjny. Zadaniem procedury jest sztuczne poprawienie rozwiązania. Procedura w pierwszym kroku zmienia pierwszą pozycję, tworząc nowe przydziały oraz określając, która ze zmian powoduje poprawę funkcji kryterium. W kolejnym kroku zmianie ulegają wartości wybranej w kroku pierwszym najlepszej kombinacji, ale zmian dokonuje się już na drugiej pozycji itd. Procedura ta trwa tak długo, aż żadna ze zmian nie poprawia czasu realizacji bądź też ustaloną na początku liczbę kroków (rys. 2).

```

wybierz ciąg
REPEAT
    FOR i:=1+p TO p+n DO
        zmień wartość i-tej współrzędnej w wybranym ciągu;
        sprawdź, czy powstała kombinacja jest lepsza;
        wyznacz najlepszą z nowych kombinacji,
        aczony przez najlepszą kombinację,
    UNTIL nie można poprawić kombinacji albo przekroczona została liczba
        iteracji

```

Rys. 2. Procedura heurystyczna

Przyjmujemy następujące oznaczenia:

$\mathbf{x}$ – rozwiązanie zakodowane w postaci ciągu złożonego z m różnych symboli długości $n+p$,

η – liczba rozwiązań – ciągów w populacji,

P^i – populacja w i -tej iteracji algorytmu $P^i = \{\mathbf{x}_1^i, \mathbf{x}_2^i, \dots, \mathbf{x}_\eta^i\}$,

$LosCh(P^i)$ – zbiór l losowo wybieranych chromosomów na potrzeby procedury heurystycznej $LosCh(P^i) = \{\mathbf{x}_{t_1}^i, \mathbf{x}_{t_2}^i, \dots, \mathbf{x}_{t_l}^i\}$,

$$\mathbf{x}_{t_j}^i \in P^i, j = 1, 2, \dots, l,$$

$Heu(\mathbf{x}, k)$ – funkcja heurystyczna poprawiająca rozwiązanie $\mathbf{x}$ (rys. 2.) przy liczbie dopuszczalnych iteracji jest równej k ,

$NewHeu(LosCh(P^i), k, l)$ – procedura, która dla każdego $\mathbf{x} \in LosCh(P^i)$, wyznacza nowe rozwiązanie $\mathbf{x}' = Heu(\mathbf{x}, k)$,

$NewPop(P^i)$ – procedura tworząca nową populację, zbiór rozwiązań za pomocą procedur ewolucji różnicowej. Z uwzględnieniem procedury ograniczonej do przydziału klientów

$$NewPop(P^i(\mathbf{x} /_{x_1, \dots, x_n}))$$

$$P^{i+1} = NewPop(P^i).$$

Hybrydowy algorytm ewolucyjny działa w sposób następujący:

0. Inicjuj dane początkowe, w sposób losowy utwórz populację $P^0 = \{\mathbf{x}_1^0, \mathbf{x}_2^0, \dots, \mathbf{x}_n^0\}$ oraz $i=0$.

1. Oceń przystosowanie populacji P^i .

2. Wygeneruj nową populację $P^{i+1} = NewPop(P^i)$ oraz zwiększ wartość i .

3. Jeśli $(i \bmod 10) = 0$, uruchom procedurę $NewHeu(LosCh(P^i), k, l)$.
4. Jeśli spełniony jest warunek stopu, otrzymane najlepsze rozwiązanie uznaj za optymalne, w przeciwnym przypadku przejdź do punktu 1.

Zaproponowane procedury zaimplementowano tworząc aplikację w Turbo Delphi pozwalającą na wprowadzanie lokalizacji klientów oraz parametrów systemu obsługi jak i również algorytmu ewolucyjnego.

3. Przykład

Zaproponowany algorytm zastosowano do znalezienia rozwiązania z udziałem 12 lokalizacji klientów, 9 potencjalnych lokalizacji systemów oraz wyborem co najwyżej 5 lokalizacji. Wyniki porównano z rozwiązaniem otrzymanym metodą przeglądu całkowitego. Wszyscy klienci posiadali takie same parametry, wszystkie centra miały te same parametry.

Na rys. 3 zaznaczono kwadratami położenie klientów, gwiazdkami lokalizację systemów obsługi. Jak również wskazano na optymalną lokalizację za pomocą linii łączących.

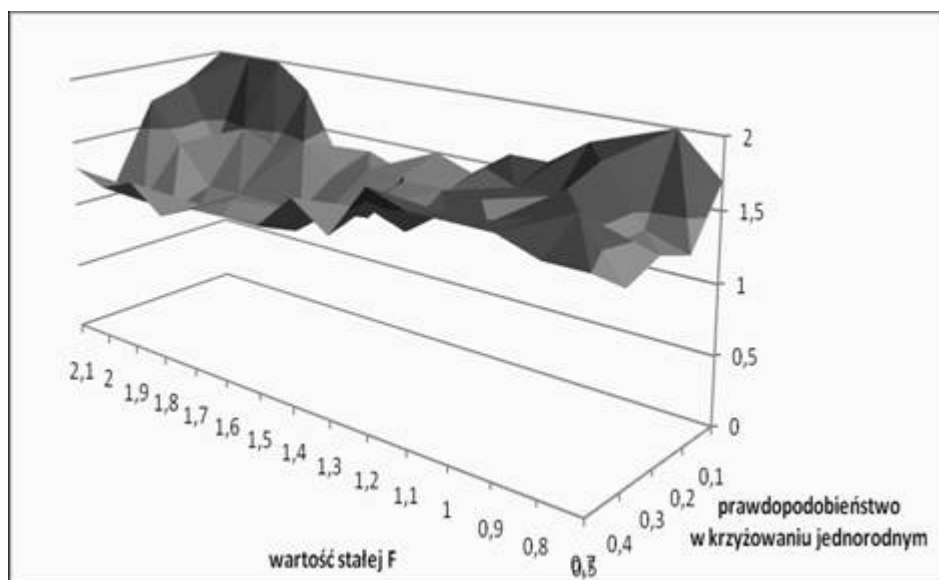


Rys. 3. Rozwiązanie problemu lokalizacyjnego

Populacja liczyła 25 rozwiązań – przykładowych lokalizacji. Dodatkowo zbadano, jak kształtuje się wpływ paramentów algorytmu: stałej skalowania F , jak i prawdopodobieństwa krzyżowania jednorodnego na efektywność algorytmu rozumianą jako stosunek numeru iteracji ze znalezionym rozwiązaniem najlepszym do najszybciej znalezionej takiego rozwiązania przy założeniu, że maksymalna liczba iteracji algorytmu – liczba analizowanych pokoleń jest równa 30 000.

Wykres 1

Zależność wpływu parametrów: stała F oraz prawdopodobieństwo krzyżowania p_c na efektywność algorytmu



Wnioski i uwagi

Algorytm przy powyższych założeniach uruchamiano dla różnych wartości jego parametrów. Z otrzymanych wyników eksperymentu wynika, że stała F powinna być równa 1.8 przy prawdopodobieństwie p_c równym 0.4. Wtedy algorytm był najefektywniejszy. Jako procedura niedeterministyczna nie zawsze znajdowane było rozwiązanie optymalne (wyznaczone dla przeglądu całkowitego). Dalsze prace nad algorytmem dotyczyć będą użycia go do rozwiązania problemu o większym rozmiarze. Algorytm może być również rozwijany ze względu na modyfikacje dotyczące sposobu kodowania rozwiązania, jak i rów-

niez samych procedur związanych z generowaniem nowych rozwiązań. Stosując metody ewolucyjne należy zwrócić uwagę na fakt, że dla rzeczywistych problemów są one alternatywą dla metody przeglądu całkowitego i metod, są również ciekawym uzupełnieniem heurystyk. Należałoby się zastanowić, czy nie zastosować heurystyk do wyznaczenia rozwiązań populacji początkowej, a następnie próbować je modyfikować za pomocą metod ewolucyjnych celem znalezienia lepszego rozwiązania.

Literatura

1. Coullard C.R., Daskin M.S., Shen Z.J. (2001). A Joint Location- Inventory Model. *Transportation Science*, 37(1), 40-55.
2. Coullard C.R., Daskin M.S., Shen Z.J. (2002). An Inventory Location Model: Formulation, Solution Algorithm and Computational Results. *Annals of Operation Research*, 110, 83-106.
3. Czachórski T. (1999). Modele kolejkowe w ocenie efektywności sieci i systemów komputerowych. Pracownia komputerowa Jacka Skalmierskiego, Gliwice.
4. Daskin M.S. (2001). A New Approach to Solving the Vertex p-center Problem to Optimality: Algorithm and Computational Results. *Communication of the Operations Research Society of Japan*, 45, 428-436.
5. Goldberg D.E. (1998). Algorytmy genetyczne i ich zastosowania. WNT, Warszawa.
6. Holland J.H. (1975). *Adaptation in Natural and Artificial Systems*. The University of Michigan Press, Ann Arbor.
7. Jakowska-Suwalska K. (2006a). Zagadnienie lokalizacji systemów masowej obsługi z wieloma klasami klientów. W: *Badania operacyjne i systemowe. Wiedza systemowa dla rozwoju regionów i przedsiębiorstw w Polsce*. Red. J. Stachowicz, A. Straszak, S. Walukiewicz. Akademicka Oficyna Wydawnicza Exit. Warszawa, 195-202.
8. Jakowska-Suwalska K. (2006b). Zagadnienie lokalizacji systemów masowej obsługi. W: *Modelowanie preferencji a ryzyko '06*. Red. T. Trzaskalik. AE, Katowice, 225-234.
9. Price K., Storm R. (1997). *Ewolucja różnicowa*. Software nr 6/97, licencja Dr. Dobb's J.
10. Price K., Storm R., Lampinen J. (2005). *Differential Evolution – A Practical Approach to Global Optimization*. Springer Verlag.
11. Sojda A. (2006). Mechanizmy ewolucyjne w zarządzaniu procesem logistycznym. *Zeszyty Naukowe. Politechnika Śląska, Gliwice*, 38, 209-219.

Joanna Utkin

Szkoła Główna Handlowa w Warszawie

KONSEKWENCJE AGREGACJI W SKOŃCZONYM MODELU TERMINOWEJ STRUKTURY STÓP PROCENTOWYCH

Wprowadzenie

Podstawową cechą modelu rynku kapitałowego jest brak możliwości arbitrażu. Celem agregacji skończonego modelu jest redukcja liczby stanów. Klaassen zaproponował metodę agregacji modelu rynku kapitałowego, która zachowuje zasadę wyceny bezarbitrażowej [2]. Dotyczyła ona skończonego rynku akcji przynoszących dywidendy, na którym obligacje występowały jako jednookresowe instrumenty bezpieczne. Metoda agregacji Klaassena jest wsteczną procedurą rekurencyjną: kolejnym agregacjom podlegają coraz dłuższe przedziały czasu o ustalonym prawym końcu. Istotą tej metody jest istnienie rekurencyjnego warunkowego prawdopodobieństwa neutralizującego ryzyko dla wydłużonych okresów. Prawdopodobieństwo to pozwala związać, poprzez równanie wyceny bezarbitrażowej, ceny na początku wydłużonego ostatniego okresu z cenami zagregowanymi na końcu okresu.

W skończonym modelu terminowej struktury stóp procentowych, określonym na sieci dwumianowej [3; 4; 5], zarówno nowe prawdopodobieństwa neutralizujące ryzyko w wydłużanych okresach, jak i nowe bezpieczne czynniki dyskontujące mają szczególną postać. Wyznamy ją agregując model metodą indukcji. Ponadto, w modelu Pedersena-Shiu-Thorlaciusa [3] zagregowane ceny końcowe mogą być przedstawione w postaci analitycznej [6]. Powyższe wyniki pozwalają wykazać prawdziwość nierówności Jarrova w podstawowych zagregowanych modelach terminowej struktury stóp procentowych, jak również proporcjonalność zdyskontowanych zysków.

1. Agregacja modelu terminowej struktury stóp procentowych

Przedmiotem agregacji jest skończony dwumianowy model terminowej struktury stóp procentowych posiadający własność Markowa [4]. W modelu występują obligacje zerokuponowe o terminach wykupu $1, \dots, M$. Dany jest horyzont agregacji T , gdzie $1 < T < M$. Jeśli agregacja jest przeprowadzona metodą Klaassena [2], to model zagregowany w przedziale $\langle 0, T \rangle$ sprowadza się do pewnego modelu jednookresowego, w którym pozostaje prawdziwa zasada wyceny bezarbitrażowej, zwana też lokalną hipotezą oczekiwań, przy czym zagregowane ceny mają dwupunktowy rozkład prawdopodobieństwa. Agregacja w przedziale $\langle 0, T \rangle$ jest osiągnięta po $T - 1$ krokach. Kolejne kroki procedury agregacji wykonuje się wstecz, agregując sukcesywnie w przedziałach $\langle T - 2, T \rangle, \langle T - 3, T \rangle, \dots, \langle 0, T \rangle$.

Metodę agregacji Klaassena w przedziale $\langle 0, T \rangle$ zastosujemy do modelu terminowej struktury kładąc nacisk na przejście indukcyjne.

Cenę obligacji zerokuponowej zapadającej w terminie m , rozpatrywaną w węźle (t, n) oznaczamy $B_t^n(m)$. Liczbę agregacji zapisujemy nad ceną obligacji. Martyngałowe prawdopodobieństwo przejścia z węzła (t, n) do węzła $(t + 1, n + 1)$ oznaczamy q_t^n .

Definicja 1

Agregację ceny obligacji (Klaassen pisze: agregacja stanu) na koniec okresu $\langle t - 1, t \rangle$ definiujemy wzorem

$$B_t^n(m) = q_{t-1}^n B_t^{n+1}(m) + (1 - q_{t-1}^n) B_t^n(m), n = 0, \dots, t - 1. \quad (1)$$

W konsekwencji równanie wyceny bezarbitrażowej obligacji na początek okresu $\langle t - 1, t \rangle$ ma postać

$$B_{t-1}^n(m) = B_{t-1}^n(t) B_t^n(m). \quad (2)$$

W szczególności cena obligacji w węźle $(T - 1, n)$ jest równa

$$B_{T-1}^n(m) = B_{T-1}^n(T) B_T^n(m).$$

Korzystając z wzorów (1) i (2) zagregowaną cenę obligacji w chwili $T-1$ (na koniec okresu $\langle T-2, T-1 \rangle$) możemy przedstawić w następujący sposób

$$\begin{aligned} B_{T-1}^{(1)}(m) &= q_{T-2}^i B_{T-1}^{i+1}(m) + (1 - q_{T-2}^i) B_{T-1}^i(m) = \\ &= q_{T-2}^i B_{T-1}^{i+1}(T) B_T^{(1)}(m) + (1 - q_{T-2}^i) B_{T-1}^i(T) B_T^{(1)}(m) = \\ &= B_{T-1}^{(1)}(T) \left[\frac{q_{T-2}^i B_{T-1}^{i+1}(T)}{B_{T-1}^i(T)} B_T^{(1)}(m) + \frac{(1 - q_{T-2}^i) B_{T-1}^i(T)}{B_{T-1}^i(T)} B_T^{(1)}(m) \right]. \end{aligned}$$

Inaczej równanie

$$B_{T-1}^{(1)}(m) = B_{T-1}^{(1)}(T) \left[a_{T-2}^i B_T^{i+1}(m) + (1 - a_{T-2}^i) B_T^i(m) \right] \quad (3)$$

ma rozwiązanie $a_{T-2}^i, a_{T-2}^i \in (0,1)$. Rozwiązanie to jest następujące

$$a_{T-2}^i = \frac{q_{T-2}^i B_{T-1}^{i+1}(T)}{B_{T-1}^i(T)}. \quad (4)$$

Związek między zagregowanymi cenami (3) pozwolił Klaassenowi określić parametr, który okaże się nowym warunkowym prawdopodobieństwem martynałowym. Ostatnim etapem pierwszej agregacji jest agregacja czasu, która „rozszerza” ostatni okres modelu na przedział $\langle T-2, T \rangle$. Stosując na przemian (2) i (3) otrzymujemy

$$\begin{aligned} B_{T-2}^i(m) &= B_{T-2}^i(T-1) B_{T-1}^{(1)}(m) = B_{T-2}^i(T-1) B_{T-1}^{(1)}(T) \\ &\left[a_{T-2}^i B_T^{i+1}(m) + (1 - a_{T-2}^i) B_T^i(m) \right] = \\ &= B_{T-2}^i(T) \left[a_{T-2}^i B_T^{i+1}(m) + (1 - a_{T-2}^i) B_T^i(m) \right]. \end{aligned}$$

Zatem prawdziwe jest nowe równanie wyceny bezarbitrażowej w okresie $\langle T-2, T \rangle$ w węźle $(T-2, i), i = 0, \dots, T-2$, mianowicie

$$B_{T-2}^i(m) = B_{T-2}^i(T) \left[a_{T-2}^i B_T^{i+1}(m) + (1 - a_{T-2}^i) B_T^i(m) \right], \quad (5)$$

w którym nowym warunkowym prawdopodobieństwem neutralizującym ryzyko jest (4).

Agregacja czasu w metodzie Klaassena doprowadza do wykluczenia transakcji w chwili $T-1$. Agregując model terminowej struktury otrzymujemy bezpieczny czynnik dyskontujący za okres $\langle T-2, T \rangle$, który jest ceną obligacji występującej na rozważanym rynku, mianowicie obligacji zerokuponowej przewidzianej do wykupu w terminie T . Na zakończenie pierwszego kroku agregacji uzyskaliśmy związek ceny obligacji długoterminowej w chwili $T-2$ i zagregowanej ceny tej obligacji w chwili T .

Jeżeli $T=2$, to procedura agregacji jest zakończona.

Wprowadzimy rekurencyjną definicję ceny k -krotnie zagregowanej, a następnie na podstawie odpowiedniego równania wyceny bezarbitrażowej wyznaczmy nowe warunkowe prawdopodobieństwo neutralizujące ryzyko dla wydłużonego okresu $\langle T-k-1, T \rangle$, gdzie $k=2, \dots, T$.

Definicja 2 (cd. definicji 1)

k -krotnie zagregowana cena obligacji jest zdefiniowana wzorem

$$B_T^{(k)}(m) = a_{T-k}^{(k-1)} B_T^{n+1}(m) + (1 - a_{T-k}^{(k-1)}) B_T^n(m), \quad n=0, \dots, T-k, \quad (6)$$

gdzie a_{T-k}^n jest liczbą z przedziału $(0,1)$, dla której spełnione jest równanie

$$B_{T-k}^n(m) = B_{T-k}^n(T) B_T^{(k)}(m). \quad (7)$$

Celem dalszych rozważań jest poszukiwanie prawdopodobieństwa neutralizującego ryzyko w wydłużonym okresie $\langle T-k-1, T \rangle$ oraz sformułowanie równania wyceny bezarbitrażowej dla tego okresu.

Korzystając z (1) i (7) rozkładamy zagregowaną cenę obligacji w chwili $T-k$

$$\begin{aligned} B_{T-k}^{(1)}(m) &= q_{T-k-1}^i B_{T-k}^{i+1}(m) + (1 - q_{T-k-1}^i) B_{T-k}^i(m) = \\ &= q_{T-k-1}^i B_{T-k}^{i+1}(T) B_T^{(k)}(m) + (1 - q_{T-k-1}^i) B_{T-k}^i(T) B_T^{(k)}(m) = \\ &= B_{T-k}^{(1)}(T) \left[\frac{q_{T-k-1}^i B_{T-k}^{i+1}(T)}{B_{T-k}^i(T)} B_T^{(k)}(m) + \frac{(1 - q_{T-k-1}^i) B_{T-k}^i(T)}{B_{T-k}^i(T)} B_T^{(k)}(m) \right]. \end{aligned}$$

Tak więc równanie wiążące zagregowane ceny na końcach okresu $\langle T-k, T \rangle$

$$B_{T-k}^{(1)}(m) = B_{T-k}^{(1)}(T) \left[a_{T-k-1}^{(k)} B_T^{i+1}(m) + (1 - a_{T-k-1}^{(k)}) B_T^i(m) \right] \quad (8)$$

ma rozwiązanie $a_{T-k-1}^i, a_{T-k-1}^i \in (0,1)$. Rozwiązanie jest równe

$$a_{T-k-1}^i = \frac{q_{T-k-1}^i B_{T-k}^{i+1}(T)}{B_{T-k}^{(1)}(T)}. \quad (9)$$

Zauważmy, że w przekształceniach prowadzących do wyznaczenia nowego prawdopodobieństwa (9) nie korzystaliśmy ze szczególnej postaci k -krotnie zagregowanej ceny obligacji, a tym samym prawdopodobieństwo (9) ma charakter uniwersalny, niezależny od agregowanej losowej wypłaty w terminie T .

Agregacja czasu rozszerza ostatni okres modelu na przedział $\langle T-k-1, T \rangle$. Stosując (2), (8) i ponownie (2) rozkładamy cenę obligacji w chwili $T-k-1$

$$\begin{aligned} B_{T-k-1}^i(m) &= B_{T-k-1}^i(T-k) B_{T-k}^{(1)}(m) = \\ &= B_{T-k-1}^{(1)}(T-k) B_{T-k}^{(1)}(T) \\ &\left[a_{T-k-1}^{(k)} B_T^{i+1}(m) + (1 - a_{T-k-1}^{(k)}) B_T^i(m) \right] = \\ &= B_{T-k-1}^i(T) \left[a_{T-k-1}^{(k)} B_T^{i+1}(m) + (1 - a_{T-k-1}^{(k)}) B_T^i(m) \right], \end{aligned}$$

skąd wynika prawdziwość nowego równania wyceny bezarbitrażowej w okresie $\langle T-k-1, T \rangle$ w węźle $(T-k-1, i)$, $i = 0, \dots, T-k-1$

$$B_{T-k-1}^i(m) = B_{T-k-1}^i(T) \left[a_{T-k-1}^{(k)} B_T^{i+1}(m) + (1 - a_{T-k-1}^{(k)}) B_T^i(m) \right]. \quad (10)$$

gdzie nowe prawdopodobieństwo neutralizujące ryzyko jest dane wzorem (9). Wzór (10) jest równaniem wyceny obligacji po k agregacjach modelu.

Zastosowanie metody agregacji Klaassena do modelu rynku obligacji wyróżnia nie tylko postać prawdopodobieństwa przejścia (9), ale także postać czynnika dyskontującego, stanowiącego w równaniu (10) cenę obligacji zapadającej w chwili T , obliczoną w momencie $T-k-1$.

Model zostaje zagregowany w okresie $\langle 0, T \rangle$ po $T-1$ agregacjach. Wówczas bezpiecznym czynnikiem dyskontującym w tym okresie jest cena obligacji $B_0(T)$, zaś prawdopodobieństwo neutralizujące ryzyko w zagregowanym modelu wynosi

$$a_0 = \frac{q_0 B_1^1(T)}{B_1^0(T)}. \quad (11)$$

Równanie wyceny bezarbitrażowej w zagregowanym modelu ma postać

$$B_0(m) = B_0(T) \left[a_0 B_T^1(m) + (1 - a_0) B_T^0(m) \right]. \quad (12)$$

Równanie wyceny (12) zostało otrzymane w wyniku podstawienia (6) do (7) w przypadku $k = T$. Ostatnia agregacja nie obejmuje czasu, ponieważ okres $\langle 0, T \rangle$, którego dotyczy równanie (12), osiągnął maksymalną długość.

Sprowadziliśmy w ten sposób model terminowej struktury stóp procentowych, określony jako model T -okresowy, do modelu jednookresowego, gdzie jedynym okresem jest przedział $\langle 0, T \rangle$. Po $T-1$ agregacjach odrzuca się możliwość obrotu obligacjami w chwilach $1, \dots, T-1$. W szczególności wyklucza się możliwość inwestowania w obligacje zapadające przed horyzontem agregacji T . Obligacja wykupywana w terminie T jest instrumentem bezpiecznym zagregowanego modelu. Obligacje w terminach wykupu $T+1, \dots, M$ są walorami ryzykownymi występującymi w zagregowanym modelu.

3. Zwroty i zdyskontowane zyski w zagregowanych modelach terminowej struktury

Definicja 3

Zwrotem z obligacji ryzykowej zapadającej w terminie m , $m = T+1, \dots, M$, w zagregowanym stanie $(T, 0)$ nazywamy liczbę

$$\bar{d}(m) = \frac{B_T^0(m)}{B_0(m)}, \quad (13)$$

a zwrotem z tej obligacji w zagregowanym stanie $(T, 1)$ liczbę

$$\bar{u}(m) = \frac{B_T^1(m)}{B_0(m)}. \quad (14)$$

Zwrotem z obligacji bezpiecznej, inaczej: czynnikiem oprocentowującym w modelu zagregowanym, nazywamy liczbę

$$\bar{w} = \frac{1}{B_0(T)}. \quad (15)$$

Warunki Jarrova w zagregowanym modelu sformułujemy analogicznie do [1]. Mają one postać podwójnych nierówności nałożonych na zwroty z obligacji o terminach wykupu $m, m = T + 1, \dots, M$, mianowicie

$$\bar{d}(m) < \bar{w} < \bar{u}(m). \quad (16)$$

Z definicji 3 wynika, że nierówność (16) można zapisać jako

$$\frac{B_T^0(m)}{B_0(m)} < \frac{1}{B_0(T)} < \frac{B_T^1(m)}{B_0(m)}. \quad (17)$$

Z równania wyceny bezarbitrażowej (12) otrzymujemy następujące przedstawienie zwrotu z obligacji bezpiecznej dla $m = T + 1, \dots, M$

$$\bar{w} = a_0 \bar{u}(m) + (1 - a_0) \bar{d}(m). \quad (18)$$

Ponieważ z (11) wynika, że $a_0 \in (0, 1)$, więc równanie (18) świadczy o tym, że dla każdego terminu zapadalności m czynnik oprocentowujący $\bar{w}$ należy do przedziału otwartego o końcach $\bar{u}(m)$ i $\bar{d}(m)$. W celu udowodnienia nierówności (17) wystarczy zatem wykazać, że dla $m = T + 1, \dots, M$ zachodzi

$$B_T^0(m) < B_T^1(m). \quad (19)$$

Podstawą badania warunków Jarrova w zagregowanych modelach terminowej struktury, niezależną od możliwości analitycznej agregacji modelu, jest twierdzenie 1.

Twierdzenie 1

Jeżeli w chwili T ceny obligacji o wszystkich terminach wykupu $m, m = T + 1, \dots, M$ spełniają nierówności

$$B_T^{i+1}(m) > B_T^i(m), i = 0, \dots, T - 1, \quad (20)$$

to po k -agregacjach metodą Klaassena, gdzie $k = 1, \dots, T - 1$, zagregowane ceny obligacji w chwili T spełniają nierówności

$$B_T^{n+1(k)}(m) > B_T^n(k)(m), n = 0, \dots, T - k - 1. \quad (21)$$

Przeprowadzimy dowód przez indukcję względem liczby agregacji.

Zgodnie z definicją 1, cena obligacji zagregowana w chwili T jest równa

$$B_T^{(1)}(m) = q_{T-1}^i B_T^{i+1}(m) + (1 - q_{T-1}^i) B_T^i(m), i = 0, \dots, T-1,$$

gdzie q_{T-1}^i jest martyngałowym prawdopodobieństwem przejścia z węzła $(T-1, i)$ do węzła $(T, i+1)$, przy czym $q_{T-1}^i \in (0, 1)$. Zagregowana cena w węźle (T, i) należy do przedziału otwartego, którego końcami są ceny niezagregowane w węzłach (T, i) i $(T, i+1)$. Z założenia (20) wynika, że cena zagregowana spełnia podwójną nierówność

$$B_T^i(m) < B_T^{(1)}(m) < B_T^{i+1}(m), i = 0, \dots, T-1. \quad (22)$$

Z nierówności (22) wynika naprzemienne położenie cen niezagregowanych i zagregowanych obligacji ryzykownych

$$B_T^0(m) < B_T^{(1)}(m) < B_T^1(m) < \dots < B_T^{T-1}(m) < B_T^{(1)}(m) < B_T^T(m).$$

Stąd otrzymujemy nierówności

$$B_T^n(m) < B_T^{(1)}(m) < B_T^{n+1}(m), n = 0, \dots, T-1. \quad (23)$$

Zakładamy, że ceny obligacji w chwili T , zagregowane $k-1$ razy, spełniają nierówności

$$B_T^{(k-1)}(m) < B_T^{(k-1)}(m) < B_T^{(k-1)}(m), n = 0, \dots, T-k. \quad (24)$$

k -krotnie zagregowana cena obligacji w chwili T , $k = 2, \dots, T-1$, jest według definicji 2 średnią ważoną cen tej obligacji zagregowanych $k-1$ razy, przy czym parametr średniej we wzorze (6) $a_{T-k}^n \in (0, 1)$. Z Założenia indukcyjnego (24) i z definicji 2 wynika więc, że

$$B_T^{(k-1)}(m) < B_T^{(k)}(m) < B_T^{(k-1)}(m), n = 0, \dots, T-k. \quad (25)$$

Prawdziwy jest zatem ciąg nierówności

$$B_T^{(k-1)}(m) < B_T^{(k)}(m) < B_T^{(k-1)}(m) < \dots < B_T^{T-k}(m) < B_T^{(k)}(m) < B_T^{T-k+1}(m),$$

skąd otrzymujemy tezę indukcyjną

$$B_T^n(m) < B_T^{(k)}(m), \quad n = 0, \dots, T-k-1, \quad (26)$$

gdzie $k = 2, \dots, T-1$.

Wniosek 1

Jeżeli w skończonym dwumianowym modelu terminowej struktury stóp procentowych ceny obligacji w chwili T spełniają założenie (20), to w modelu zagregowanym w okresie $\langle 0, T \rangle$ zachodzi nierówność (19), a w konsekwencji spełnione są warunki Jarrova (16).

Wniosek 2

Model Sandmanna-Sondermanna [5] dopasowany do początkowych cen obligacji spełniających nierówności

$$0 < B_0(M) < \dots < B_0(1) < 1, \quad (27)$$

po agregacji w okresie $\langle 0, T \rangle$, gdzie $1 < T < M$, spełnia warunki Jarrova (16).

Wniosek 3

Model Pedersena-Shiu-Thorlaciusa [3] dopasowany do początkowych cen obligacji spełniających nierówności (27), w którym zmienności cen obligacji jednookresowych $c_t = B_t^{i+1}(t+1)/B_t^i(t+1)$, $i = 0, \dots, t-1$, spełniają podwójną nierówność

$$1 < c_t < \sqrt[t]{\frac{B_0(t)}{B_0(t+1)}}, \quad t = 1, \dots, M-1, \quad (28)$$

po agregacji w okresie $\langle 0, T \rangle$, gdzie $1 < T < M$, spełnia warunki Jarrova (16).

Definicja 4

Zdyskontowanym zyskiem z obligacji w modelu zagregowanym w okresie $\langle 0, T \rangle$ nazywamy zmienną losową

$$\bar{z}_m = B_0(T) B_T^{(T-1)}(m) - B_0(m),$$

gdzie m jest terminem wykupu obligacji, $m = T+1, \dots, M$.

W zagregowanym stanie (T, i) , $i = 0, 1$ zdyskontowany zysk ma wartość

$$\bar{z}_m^i = B_0(T) B_T^{(T-1)}(m) - B_0(m). \quad (29)$$

Wyrażając $B_0(m)$ za pomocą wyceny bezarbitrażowej, zysk (29) przedstawiamy w postaci

$$\bar{z}_m^i = B_0(T) [B_T^{(T-1)}(m) - a_0 B_T^1(m) - (1 - a_0) B_T^0(m)], \quad i = 0, 1.$$

W poszczególnych stanach otrzymujemy

$$\begin{aligned}\bar{z}_m^0 &= -a_0 B_0(T) (B_T^{(T-1)}(m) - B_T^{(T-1)}(m)), \\ \bar{z}_m^1 &= (1 - a_0) B_0(T) (B_T^{(T-1)}(m) - B_T^{(T-1)}(m)).\end{aligned}$$

Jeżeli w zagregowanym modelu zachodzi nierówność (19), to możemy obliczyć iloraz zdyskontowanych zysków w zagregowanych stanach, mianowicie

$$\frac{\bar{z}_m^1}{\bar{z}_m^0} = -\frac{1 - a_0}{a_0}.$$

Prawdziwe jest więc twierdzenie 2.

Twierdzenie 2

Jeżeli w zagregowanym modelu terminowej struktury stóp procentowych w okresie $\langle 0, T \rangle$ zachodzą nierówności

$$B_T^{(T-1)}(m) < B_T^{(T-1)}(m), m = T + 1, \dots, M,$$

to zdyskontowane zyski z obligacji ryzykownych w zagregowanych stanach $(T, 0)$ i $(T, 1)$ są proporcjonalne

$$\frac{\bar{z}_m^1}{\bar{z}_m^0} = -\frac{1 - a_0}{a_0}, \quad m = T + 1, \dots, M. \quad (30)$$

3. Prawdopodobieństwo subiektywne w modelu zagregowanym

Prawdopodobieństwo subiektywne, zwane też prawdopodobieństwem rzeczywistym, nie jest wykorzystywane przy konstrukcji skończonego modelu terminowej struktury stóp procentowych. Podobnie dzieje się w przypadku modelu zagregowanego metodą Klaassena. Autor metody agregacji dopuszcza różne alternatywne sposoby określenia zagregowanego prawdopodobieństwa subiektywnego pozostawiając problem otwarty. Specyfika skończonych modeli terminowej struktury okazuje się pomocna przy rozwiązywaniu tego problemu. Mianowicie w zagregowanym modelu dwumianowym występują dwa stany końcowe. Jeśli w modelu przed agregacją wyróżniony jest kierunek hossy, to ten sam kierunek hossy występuje po agregacji. Co więcej, gdy zagregowane ceny obligacji ryzykownych w jednym zagregowanym stanie wzrosły, a w dru-

gim – spadły w porównaniu z ceną wykupu obligacji bezpiecznej, to wzrost wszystkich cen obligacji ryzykownych można przypisać pierwszemu stanowi, a spadek wszystkich cen – drugiemu.

Wniosek 4

Jeżeli w zagregowanym modelu terminowej struktury stóp procentowych spełnione są warunki Jarrowa (16), to w zagregowanym stanie $(T,1)$ następuje wzrost cen wszystkich obligacji ryzykownych, a w zagregowanym stanie $(T,0)$ – spadek.

Oznaczając przez $\bar{P}$ funkcję prawdopodobieństwa na zbiorze zagregowanych stanów końcowych $\{(T,0), (T,1)\}$, przy spełnionych warunkach Jarrowa, wartość $\bar{P}((T,1))$ interpretujemy jako subiektywne prawdopodobieństwo wzrostu cen, zaś $\bar{P}((T,0))$ jako subiektywne prawdopodobieństwo spadku w modelu zagregowanym.

Jeśli znamy subiektywne prawdopodobieństwo wzrostu cen obligacji długoterminowych w okresie $\langle 0, T \rangle$, to stosując model, w którym zachodzą warunki Jarrowa, przyjmujemy wartość $\bar{P}((T,1))$ równą danemu prawdopodobieństwu.

Zaletą agregacji skończonego modelu terminowej struktury stóp procentowych jest redukcja liczby stanów w okresie $\langle 0, T \rangle$ do dwóch, z których jeden charakteryzuje się wzrostem, a drugi spadkiem cen. Wówczas szacowanie subiektywnego prawdopodobieństwa wzrostu albo spadku cen z upływem T okresów może być efektywnie przeprowadzone. W szczególności może być wykorzystana znajomość innego rozkładu cen obligacji.

4. Zobowiązania w zagregowanym modelu terminowej struktury

Metoda agregacji została wprowadzona przez Klaassena dla akcji z dywidendami [2]. Ze względu na specyfikę modeli terminowej struktury stóp procentowych, wyodrębniliśmy agregację obligacji. W trakcie agregacji obligacji wyznaczyliśmy warunkowe prawdopodobieństwa martyngałowe oraz bezpieczne czynniki dyskontujące, które występują w równaniu wyceny bezarbitrażowej w wydłużonym okresie. Dysponując wymienionymi parametrami agregacji rozszerzymy definicje agregacji cen 1 i 2 na losowe zobowiązania płatne w terminie T . Wartość zobowiązania w danym węźle ma postać liniowej funkcji cen obligacji w tym węźle o współczynnikach zależnych od węzła.

Rozważamy zobowiązanie L_T o wartościach L_T^i danych w węzłach $(T, i), i = 0, \dots, T$.

Definicja 5

Agregację zobowiązania L_T w węźle (T, i) definiujemy wzorem (1)

$$L_T^{(1)} = q_{T-1}^i L_T^{i+1} + (1 - q_{T-1}^i) L_T^i, \quad i = 0, \dots, T-1. \quad (31)$$

Wycena zobowiązania w chwili $T-1$ ma wtedy postać (2)

$$L_{T-1}^i = B_{T-1}^i(T) L_T^{(1)}, \quad i = 0, \dots, T-1. \quad (32)$$

Definicja 6 (cd. definicji 5)

k -krotną agregację zobowiązania w węźle (T, i) , przy czym $k = 2, \dots, T$, definiujemy wzorem rekurencyjnym (6)

$$L_T^{(k)} = a_{T-k}^i L_T^{(k-1)} + (1 - a_{T-k}^i) L_T^{(k-1)}, \quad i = 0, \dots, T-k, \quad (33)$$

gdzie a_{T-k}^i jest rozwiązaniem równania wyceny na moment $T-k$ (7)

$$L_{T-k}^i = B_{T-k}^i(T) L_T^{(k)}, \quad i = 0, \dots, T-k. \quad (34)$$

Rozwiązanie ma postać (9)

$$a_{T-k}^i = \frac{q_{T-k}^i B_{T-k+1}^{i+1}(T)}{B_{T-k+1}^i(T)}. \quad (35)$$

Model dwumianowy jest zagregowany w okresie $\langle 0, T \rangle$ po $T-1$ agregacjach. W modelu tym zmienna losowa L_T ma dwupunktowy rozkład prawdopodobieństwa. W wyniku T -krotnej agregacji ceny otrzymujemy wartość początkową zobowiązania równą

$$L_0 = B_0(T) [a_0 L_T^1 + (1 - a_0) L_T^0]. \quad (36)$$

Równanie (36) przedstawia wycenę bezarbitrażową zobowiązania L_T w okresie $\langle 0, T \rangle$ w modelu zagregowanym $T-1$ razy.

5. Minimalizacja straty w zagregowanym modelu terminowej struktury

Dane jest losowe zobowiązanie płatne w chwili T, L_T oraz kapitał początkowy v . Wycena bezarbitrażowa na moment 0 losowego zobowiązania jest dana wzorem (36).

Jeżeli $v < L_0$, to wystąpi losowy niedobór. W minimalizacji niedoboru wykorzystamy model terminowej struktury zagregowany w okresie $\langle 0, T \rangle$. Zdyskontowana losowa strata jest zdefiniowana następująco

$$-\bar{X} = B_0(T) \left(L^{(T-1)} - V_T^{(T-1)} \right). \quad (37)$$

Interesuje nas rozwiązanie następującego problemu minimalizacji oczekiwanej straty

$$\min E_{\bar{P}}(-\bar{X})^+ \quad (38)$$

pod warunkiem $V_0 = v$.

W modelu terminowej struktury zagregowanym w okresie $\langle 0, T \rangle$ strategie mają postać $\varphi = (0, \dots, 0, \varphi_T, \dots, \varphi_M)$. Warunek z zadania (38) sprowadza się do równania

$$\varphi_T B_0(T) + \sum_{m=T+1}^M \varphi_m B_0(m) = v.$$

Stąd wynika, że zdyskontowana zagregowana końcowa wartość portfela jest po eliminacji φ_T równa

$$B_0(T) V_T^{(T-1)} = v + \sum_{m=T+1}^M \varphi_m (B_0(T) B_T^{(T-1)}(m) - B_0(m)). \quad (39)$$

Biorąc pod uwagę definicję zdyskontowanego zysku w modelu zagregowanym, miarę straty przedstawimy jako

$$\begin{aligned} E_{\bar{P}}(-\bar{X})^+ &= \bar{P} \left(B_0(T) L_T^{(T-1)} - v - \sum_{m=T+1}^M \varphi_m \bar{Z}_m^1 \right)^+ + \\ &+ (1 - \bar{P}) \left(B_0(T) L_T^0 - v - \sum_{m=T+1}^M \varphi_m \bar{Z}_m^0 \right)^+. \end{aligned} \quad (40)$$

Na podstawie twierdzenia 2 o proporcjonalności zdyskontowanych zysków sprowadzamy oczekiwaną stratę (40) do postaci

$$\begin{aligned} E_{\bar{p}}(-\bar{X})^+ &= \bar{p} \left(B_0(T) L_T^{(T-1)} - v + \frac{1-a_0}{a_0} \sum_{m=T+1}^M \varphi_m \bar{z}_m^0 \right)^+ + \\ &+ (1-\bar{p}) \left(B_0(T) L_T^{(T-1)} - v - \sum_{m=T+1}^M \varphi_m \bar{z}_m^0 \right)^+. \end{aligned} \quad (41)$$

Dla strategii φ minimalizującej (41) przynajmniej jeden ze składników jest równy zero. Jeżeli

$$\sum_{m=T+1}^M \varphi_m \bar{z}_m^0 = B_0(T) L_T^{(T-1)} - v,$$

to

$$B_0(T) L_T^{(T-1)} - v + \frac{1-a_0}{a_0} \left(B_0(T) L_T^{(T-1)} - v \right) = \frac{1}{a_0} (L_0 - v).$$

Jeżeli

$$\sum_{m=T+1}^M \varphi_m \bar{z}_m^0 = \frac{a_0}{a_0 - 1} \left(B_0(T) L_T^{(T-1)} - v \right),$$

to

$$B_0(T) L_T^{(T-1)} - v + \frac{a_0}{1-a_0} \left(B_0(T) L_T^{(T-1)} - v \right) = \frac{1}{1-a_0} (L_0 - v).$$

W rezultacie najmniejsza wartość oczekiwanej straty (41) jest równa

$$\sigma_{\min} = (L_0 - v) \min \left\{ \frac{\bar{p}}{a_0}, \frac{1-\bar{p}}{1-a_0} \right\}. \quad (42)$$

Następny etap optymalizacji polega na wyznaczeniu strategii φ , stałych w $\langle 0, T \rangle$, dla których funkcja straty osiąga najmniejszą wartość (42).

Przykład

Rozważamy terminową strukturę Ho-Lee [3] dla $M = 4, q = \frac{1}{2}$, zagregowaną w okresie $\langle 0, 2 \rangle$ ze stałą s spełniającą założenie (28) z wniosku 3. Przyjmujemy prawdopodobieństwo osiągnięcia węzła (2,1) w modelu zagregowanym równe $\frac{2}{5}$. Zakładamy, że zobowiązanie ma wartości $L_2^2 = 20, L_2^1 = 10, L_2^0 = 10$,

zaś kapitał zainwestowany w portfel wynosi $10B_0(2)$. Wyznamy najmniejszą wartość oczekiwaną straty, a następnie wyznaczmy strategię, dla której ta wartość jest przyjęta.

Na podstawie definicji 5 obliczamy $L_2^{(1)} = 15, L_2^{(0)} = 10$. Nowe prawdopodobieństwo martyngałowe dla okresu $\langle 0, 2 \rangle$ jest równe $a_0 = \frac{1}{1+s}$.

Z (36) wynika, że wycena zobowiązania na chwilę zero wynosi

$$L_0 = B_0(2) \left(a_0 L_2^{(1)} + (1 - a_0) L_2^{(0)} \right) = B_0(2) \frac{15 + 10s}{1 + s},$$

stąd

$$L_0 - v = B_0(2) \frac{5}{1 + s}.$$

Ponieważ $L_0 - v > 0$, zajmiemy się minimalizacją miary straty.

Obliczamy gęstości cen zagregowanych stanów

$$\bar{p} = \frac{2}{5}, \text{ więc } \frac{\bar{p}}{a_0} = \frac{2(1+s)}{5}, \frac{1-\bar{p}}{1-a_0} = \frac{3(1+s)}{5s}.$$

Dla $s \in (0, 1)$ zachodzi $2s < 3$, więc zgodnie z (42)

$$\sigma_{\min} = 2B_0(2). \quad (43)$$

Wartość (43) jest przyjęta, gdy drugi składnik (41) jest równy 0. Stąd ilości obligacji ryzykownych φ_3 i φ_4 spełniają równanie

$$\varphi_3 \bar{z}_3^0 + \varphi_4 \bar{z}_4^0 = 0.$$

Według definicji 4 zdyskontowane zyski z obligacji ryzykownych w stanie 0 są dane wzorami

$$\bar{z}_3^0 = B_0(2) B_2^{(0)}(3) - B_0(3),$$

$$\bar{z}_4^0 = B_0(2) B_2^{(0)}(4) - B_0(4).$$

Wykorzystamy ceny obligacji w chwili 2 wyznaczone dla modeli Ho-Lee ze stałą S i z warunkowym prawdopodobieństwem martyngałowym równym 0,5, mianowicie

$$B_2^2(3) = \frac{B_0(3)}{B_0(2)} \frac{2}{1+s^2}, \quad B_2^1(3) = B_2^2(3)s, \quad B_2^0(3) = B_2^2(3)s^2,$$

$$B_2^2(4) = \frac{B_0(4)}{B_0(2)} \frac{2}{(1+s^2)(1-s+s^2)}, \quad B_2^1(4) = B_2^2(4)s^2, \quad B_2^0(4) = B_2^2(4)s^4.$$

Z definicji 1 wynika, że

$$B_2^{(1)}(3) = \frac{B_0(3)}{B_0(2)} \frac{s + s^2}{1 + s^2}, \quad B_2^{(1)}(4) = \frac{B_0(4)}{B_0(2)} \frac{s^2}{1 - s + s^2},$$

stąd otrzymujemy zyski z obligacji równe

$$\bar{z}_3^0 = B_0(3) \frac{s - 1}{1 + s^2},$$

$$\bar{z}_4^0 = B_0(4) \frac{s - 1}{1 - s + s^2}.$$

Zatem związek ilości obligacji ryzykownych w portfelu jest następujący

$$\varphi_3 = -\frac{B_0(4)}{B_0(3)} \frac{1 + s^2}{1 - s + s^2} \varphi_4, \quad \varphi_4 \in R. \quad (44)$$

Z ustalenia początkowej wartości portfela równej $10B_0(2)$ wyznaczamy ilość obligacji bezpiecznej φ_2

$$\varphi_2 B_0(2) + \varphi_3 B_0(3) + \varphi_4 B_0(4) = 10B_0(2).$$

Stąd po uwzględnieniu (44) otrzymujemy

$$\varphi_2 = 10 + \frac{B_0(4)}{B_0(2)} \frac{s}{1 - s + s^2} \varphi_4, \quad \varphi_4 \in R. \quad (45)$$

Strategie $\varphi = (0, \varphi_2, \varphi_3, \varphi_4)$, w których φ_2, φ_3 są liniowymi funkcjami zmiennej $\varphi_4, \varphi_4 \in R$, generują minimalną stratę (43).

Powyższy przykład dotyczył minimalizacji oczekiwanej straty w wybranym modelu dwumianowej struktury terminowej. Dzięki agregacji została wykorzystana strata w ujęciu statycznym. Podobnie można stosować inne statyczne miary straty, oparte na danych funkcjach charakteryzujących stosunek inwestora do straty. Można również modelować, statyczne z natury, miary ryzyka na płaszczyźnie. Zagadnieniem typowym dla aplikacji terminowej struktury jest analiza zobowiązań w danym horyzoncie czasu – takie problemy mogą być modelowane na gruncie statycznym dzięki znajomości metody agregacji struktury terminowej.

Literatura

1. Jarrow R. (2002). Modelling Fixed Income Securities and Interest Rate Options. Stanford Univ. Press, Stanford.
2. Klaassen P. (1998). Financial Asset Pricing Theory and Stochastic Programming Models for Asset Liability Management. A Synthesis. Management Science, 44, 31-48.

-
3. Pedersen H., Shiu E., Thorlacius A. (1989). Arbitrage Free Pricing of Interest Rate Contingent Claims. *Transactions of the Society of Actuaries*, 41, 237-279.
 4. Pliska S. (2005). Wprowadzenie do matematyki finansowej. Modele z czasem dyskretnym. WNT, Warszawa.
 5. Sandmann K., Sondermann D. (1993). A Term Structure Model and the Pricing of Interest Rate Derivatives. *Review of Future Markets*, 12, 391-423.
 6. Utkin J. (2009). Aggregation Analysis in Pedersen-Shiu-Thorlacius Model. SGH, Warszawa.

Sławomir Żułka

Europejski Bank Inwestycyjny w Luksemburgu

DYNAMICZNA OPTYMALIZACJA STRATEGII INWESTYCYJNEJ W MODELU HO-LEE

Wstęp

Opracowanie stanowi kontynuację i rozszerzenie na przypadki dynamiczne badań wybranych problemów przedstawionych w pracy [4] i dotyczących zabezpieczenia spłaty zobowiązania w dyskretnych i skończonych modelach ewolucji terminowej struktury stóp procentowych. Celem pracy jest prezentacja wyników poszukiwania optymalnych, dynamicznych strategii inwestycyjnych w dyskretnym i skończonym modelu stóp procentowych. Optymalna, dynamiczna strategia inwestycyjna definiowana jest jako sekwencja jednookresowych i warunkowanych aktualnym stanem modelu decyzji alokacji majątku pomiędzy obciążone ryzykiem aktywa inwestycyjne w sposób minimalizujący funkcję straty, będącą miarą niedotrzymania płatnego w przyszłości zobowiązania. Problem, który formułujemy, polega na odpowiedzi na pytanie, czy możliwe jest znalezienie takiej strategii inwestycyjnej, która przy zadanej początkowej wartości majątku redukuje do zera wartość oczekiwaną przyjętej funkcji straty. Modelem ewolucji stóp procentowych rozważanym w opracowaniu jest specyfikacja zaproponowana przez Ho-Lee [2], natomiast minimalizowana funkcja straty zapożyczona została od Cvitanica i Karatzasa [1].

1. Specyfikacja modelu Ho-Lee

Model Ho-Lee jest skończonym, dwumianowym modelem terminowej struktury stóp procentowych. W modelu tym ewolucja struktury stóp procentowych opisana jest za pomocą rekombinującego się drzewa, w którym ceny obligacji zależą od czasu oraz liczby wzrostów stóp, ale nie zależą od sekwencji występowania tychże wzrostów.

Specyfikacja modelu Ho-Lee zakłada znajomość zestawu początkowych cen obligacji zerokuponowych, zapadających w momentach od 1 do M . Oznaczając przez $B_t^n(m)$ cenę obligacji zerokuponowej zapadającej w momencie m , a obserwowanej w momencie t i przy warunku zajścia n wzrostów stóp do momentu t , odpowiada to znajomości wektora cen $B_0^0(1)$, $B_0^0(2)$, ..., $B_0^0(M)$, na podstawie których można jednoznacznie wyznaczyć odpowiadającą im krzywą stóp zerokuponowych $r_0^0(1)$, $r_0^0(2)$, ..., $r_0^0(M)$. Model Ho-Lee wymaga również przyjęcia współczynnika zmienności stóp procentowych o postaci $\frac{1}{k}$, gdzie k jest liczbą rzeczywistą większą od jedności, dobraną w taki sposób, aby proces cen generowany przez model nie prowadził do występowania ujemnych stóp procentowych (w praktyce oznacza to, że k nie może nadmiernie odbiegać od jedności). Współczynnik zmienności stóp procentowych w modelu Ho-Lee jest ustalony, jednakże w uogólnionych modelach może być wyrażony jako funkcja czasu i stanu. Proces rządzący ewolucją stóp procentowych wyspecyfikowany jest w taki sposób, aby spełnione było następujące równanie

$$B_t^{n+1}(t+1) = \frac{1}{k} B_t^n(t+1) \quad (1)$$

a ceny obligacji w modelu spełniały warunek wyceny bezarbitrażowej

$$B_t^n(m) = B_t^n(t+1)[qB_{t+1}^{n+1}(m) + (1-q)B_{t+1}^n(m)] \quad (2)$$

gdzie q oznacza niezależne od czasu i dotychczasowej liczby wzrostów stóp prawdopodobieństwo wystąpienia wzrostu stóp w następnym okresie. Choć q , podobnie jak współczynniki zmienności stóp procentowych, może być wyrażone w sposób ogólny jako funkcja czasu lub stanu, w oryginalnym modelu Ho-Lee ustalone jest ono na poziomie 0,5. Warunki (1) i (2), przy znajomości początkowej terminowej struktury stóp zerokuponowych, służą do wyznaczenia drzewa ewolucji cen (co jest równoznaczne z wyznaczeniem drzewa ewolucji stóp rentowności) dla każdej z obligacji, które obserwowane są w momencie zerowym, aż do terminu ich zapadalności.

Można wykazać (por. aneks), że w momencie t , w stanie odpowiadającym n wzrostom stóp procentowych w przeszłości, cena obligacji zapadającej w momencie m wyrażona jest wzorem

$$B_t^n(m) = \frac{B_0^0(m)}{B_0^0(t)} (k^{(t-1)-(m-1)})^n \prod_{j=0}^{t-1} \frac{g[j-(t-1)]}{g[j-(m-1)]} \quad (3)$$

gdzie g jest funkcją pomocniczą o postaci

$$g(j-s) = [1-q + qk^{j-s}] \quad (4)$$

2. Sformułowanie problemu zabezpieczenia spłaty zobowiązania

Rozważamy model rynku, na którym w momencie początkowym dostępne są cztery obligacje zerokuponowe, zapadające odpowiednio w momencie pierwszym, drugim, trzecim i czwartym. Zgodnie z przyjętymi oznaczeniami, ich ceny wyrażamy przez $B_0^0(1)$, $B_0^0(2)$, $B_0^0(3)$ oraz $B_0^0(4)$. Przyjmujemy dwuokresowy horyzont inwestycyjny, co oznacza, że zobowiązanie L płatne jest w momencie drugim (na końcu okresu drugiego). W momencie pierwszym obligacja o najkrótszym terminie zapadalności zapada, natomiast ceny pozostałych obligacji przyjmują wartości zależne od tego, czy model znalazł się w stanie odpowiadającym wzrostowi stop procentowych ($n=1$) czy w stanie odpowiadającym spadkowi stop procentowych ($n=0$). W momencie drugim obligacja o terminie zapadalności w momencie dwa zapada, a obligacje o terminie zapadalności trzy i cztery przyjmują nowe ceny, zależne od tego, czy model znalazł się w stanie n równym zero, jeden lub dwa. W konsekwencji przyjętej wartości prawdopodobieństwa zmiany stanu równej $q=0,5$, pozostanie modelu w stanie n oraz przejście do stanu $n+1$ jest zawsze jednakowo prawdopodobne. Zgodnie ze wzorem (3), ewolucje cen obligacji o zapadalnościach w momentach dwa, trzy i cztery opisują poniższe wzory

$$\left\{ \begin{array}{l} B_1^1(2) = \frac{B_0^0(2)}{B_0^0(1)} \frac{2}{k+1} \\ B_1^0(2) = \frac{B_0^0(2)}{B_0^0(1)} \frac{2k}{k+1} \end{array} \right\}, \left\{ \begin{array}{l} B_1^1(3) = \frac{B_0^0(3)}{B_0^0(1)} \frac{2}{k^2+1} \\ B_1^0(3) = \frac{B_0^0(3)}{B_0^0(1)} \frac{2k^2}{k^2+1} \end{array} \right\}, \left\{ \begin{array}{l} B_1^1(4) = \frac{B_0^0(4)}{B_0^0(1)} \frac{k^{-3}}{0.5+0.5k^{-3}} \\ B_1^0(4) = \frac{B_0^0(4)}{B_0^0(1)} \frac{1}{0.5+0.5k^{-3}} \end{array} \right\}$$

$$\left\{ \begin{array}{l} B_2^2(3) = \frac{B_0^0(3)}{B_0^0(2)} \frac{k^{-2}}{0.5+0.5k^{-2}} \\ B_2^1(3) = \frac{B_0^0(3)}{B_0^0(2)} \frac{k^{-1}}{0.5+0.5k^{-2}} \\ B_2^0(3) = \frac{B_0^0(3)}{B_0^0(2)} \frac{1}{0.5+0.5k^{-2}} \end{array} \right\}, \left\{ \begin{array}{l} B_2^2(4) = \frac{B_0^0(4)}{B_0^0(2)} \frac{k^{-4}(0.5+0.5k^{-1})}{(0.5+0.5k^{-3})(0.5+0.5k^{-2})} \\ B_2^1(4) = \frac{B_0^0(4)}{B_0^0(2)} \frac{k^{-2}(0.5+0.5k^{-1})}{(0.5+0.5k^{-3})(0.5+0.5k^{-2})} \\ B_2^0(4) = \frac{B_0^0(4)}{B_0^0(2)} \frac{(0.5+0.5k^{-1})}{(0.5+0.5k^{-3})(0.5+0.5k^{-2})} \end{array} \right\}$$

Przyjmujemy, że inwestor dysponuje majątkiem początkowym o wielkości $v = V_0$. Majątek ten posłuży do spłacenia zapadającego w przyszłości zobowiązania o wielkości L , przy czym zdyskontowana do momentu zerowego wartość tego majątku równa jest początkowej wielkości majątku inwestora:

$$B_0(T)L = V_0 \quad (5)$$

Stopień niedotrzymania zobowiązania w każdym ze stanów w terminie horyzontu mierzony jest za pomocą funkcji straty Cvitanica i Karatzasa o postaci

$$S = B_0(T)(L - V_T)^+ \quad (6)$$

gdzie S jest wartością funkcji straty, a $(x)^+ = \max\{x; 0\}$. W przypadku, gdy końcowa wartość majątku przekracza wartość zobowiązania, wyrażenie w nawiasie jest ujemne i funkcja straty zwraca wartość zero. W przeciwnym wypadku, funkcja straty przyjmuje wartość równą zdyskontowanej do momentu zero wielkości niedoboru majątku w stosunku do wielkości zobowiązania.

Inwestor stoi przed problemem wyboru dynamicznej strategii inwestycyjnej gwarantującej zwrot z zainwestowanego majątku w takiej wysokości, która zminimalizuje do zera wartość oczekiwaną z funkcji straty. W rozważaniach wykluczamy możliwość dokonywania krótkiej sprzedaży oraz pożyczania pieniędzy na rynku, natomiast możliwość dokonywania transakcji (bez kosztów transakcyjnych) dopuszczamy jedynie w momentach zero i jeden. Wybór strategii sprowadza się więc do znalezienia trzech wektorów alokacji majątku pomiędzy dostępne aktywa inwestycyjne:

- wektora udziałów czterech obligacji dostępnych w momencie zerowym,
- wektora udziałów trzech obligacji dostępnych w momencie jeden, obowiązującego w przypadku, gdy stopy procentowe wzrosły (w stanie $n=1$),
- wektora udziałów trzech obligacji dostępnych w momencie jeden, obowiązującego w przypadku, gdy stopy procentowe spadły (w stanie $n=0$).

Udziały te inwestor musi dobrać w taki sposób, aby wartość oczekiwana funkcji straty w terminie horyzontu wynosiła zero

$$E[B_0(T)(L - V_T)^+] = 0 \quad (7)$$

Ponieważ prawdopodobieństwa wystąpienia wszystkich stanów są dodatnie, a funkcja straty nie przyjmuje wartości ujemnych, wartość oczekiwana funkcji straty przyjmie wartość zerową tylko wtedy, gdy wszystkie możliwe wartości, jakie ta zmienna losowa może przyjąć, będą zerowe. Stanie się tak tylko wtedy, gdy w każdym stanie modelu w terminie horyzontu wartość końcowa majątku będzie przynajmniej równa wartości zobowiązania. Ze względu na (5), wymóg ten spełniony jest w sposób trywialny w sytuacji, gdy majątek początkowy inwestora zainwestowany jest wyłącznie w obligację o terminie zapadalności dwa lub w portfel obligacji replikujący syntetycznie zwroty z tejże obligacji w każdym okresie i stanie modelu. Pojawia się pytanie, czy istnieje strategia inwestycyjna inna niż powyżej opisana (tzn. inna niż replikująca zwroty z obligacji o terminie zapadalności dwa), która gwarantuje zerowanie się wartości oczekiwanej ze zdefiniowanej wzorem (6) funkcji straty?

3. Rezultat badania i jego interpretacja

Poszukując odpowiedzi na sformułowane powyżej pytanie wygodnie jest rozpocząć od wyznaczenia zwrotów ze strategii inwestycyjnej dla każdej ze ścieżek drzewa ewolucji cen w modelu. Oznaczając przez $x_t^n(m)$ obowiązujący w momencie t i stanie n udział w portfelu inwestora obligacji zapadającej w chwili m , możemy przedstawić końcową wartość majątku inwestora dla każdej ze ścieżek ewolucji cen w modelu jako iloczyn zwrotu ze strategii w pierwszym i w drugim okresie. Przy takich oznaczeniach, końcowe wartości majątku inwestora wynoszą odpowiednio:

- dla cen ze stanu 1 w pierwszym okresie i 2 w drugim okresie

$$V_T = V_0 \frac{1}{B_0^0(2)} \left(x_0(1) + x_0(2) \frac{2}{(k+1)} + x_0(3) \frac{2}{(k^2+1)} + x_0(4) \frac{2}{(k^3+1)} \right) \cdot \left(x_1^1(2) \frac{k+1}{2} + x_1^1(3) + x_1^1(4) \frac{1+k}{k^2+1} \right)$$

- dla cen ze stanu 1 w pierwszym okresie i 1 w drugim okresie

$$V_T = V_0 \frac{1}{B_0^0(2)} \left(x_0(1) + x_0(2) \frac{2}{(k+1)} + x_0(3) \frac{2}{(k^2+1)} + x_0(4) \frac{2}{(k^3+1)} \right) \cdot \left(x_1^1(2) \frac{k+1}{2} + x_1^1(3)k + x_1^1(4) \frac{k^2+k^3}{k^2+1} \right)$$

- dla cen ze stanu 0 w pierwszym okresie i 1 w drugim okresie

$$V_T = V_0 \frac{1}{B_0^0(2)} \left(x_0(1) + x_0(2) \frac{2k}{(k+1)} + x_0(3) \frac{2k^2}{(k^2+1)} + x_0(4) \frac{2k^3}{(k^3+1)} \right) \cdot \left(x_1^0(2) \frac{k+1}{2k} + x_1^0(3)k^{-1} + x_1^0(4) \frac{1+k^{-1}}{k^2+1} \right)$$

- natomiast dla cen ze stanu 0 zarówno w pierwszym, jak i drugim okresie

$$V_T = V_0 \frac{1}{B_0^0(2)} \left(x_0(1) + x_0(2) \frac{2k}{(k+1)} + x_0(3) \frac{2k^2}{(k^2+1)} + x_0(4) \frac{2k^3}{(k^3+1)} \right) \cdot \left(x_1^0(2) \frac{k+1}{2k} + x_1^0(3) + x_1^0(4) \frac{k^2+k}{k^2+1} \right)$$

Aby zmienna losowa S wyrażona wzorem (6) przyjmowała zawsze wartość zero, udziały $x_t^n(m)$ muszą być tak dobrane, aby iloczyny powyższych wyrażeń w nawiasach przyjmowały wartość większą lub równą jeden we wszystkich czterech przypadkach odpowiadających czterem możliwym ścieżkom na drzewie ewolucji cen w rozważanym modelu. Dowiedzimy nie wprost, że taka strategia, przy przyjętych w modelu założeniach, nie istnieje.

Oznaczmy przez X dowolną strategię polegającą na replikowaniu wypłat obligacji zapadającej w momencie dwa; wiemy, że taka strategia w sposób trywialny spełnia wymóg zerowania wartości oczekiwanej z funkcji straty. Założymy, że istnieje strategia Y charakteryzująca się innymi wypłatami niż obligacja zapadająca w terminie dwa i również gwarantująca zerowanie się wartości oczekiwanej z funkcji straty. W sposób oczywisty podział majątku początkowego v na dwie części i jednocześnie zainwestowanie ich w strategię X i Y również musiałby zaowocować wyzerowaniem się oczekiwanej straty. W szczególności podział majątku może zostać dokonany w taki sposób, aby

dowolnie mała, ale niezerowa część majątku inwestowana była w strategię Y. Oznaczmy taką kombinowaną strategię przez X'. Strategia X w okresie pierwszym generuje odpowiednio dla stanu 1 i stanu 0 zwroty o wysokościach odpowiednio $\frac{2}{B_0^0(1)(k+1)}$ oraz $\frac{2k}{B_0^0(1)(k+1)}$. Strategia X' z definicji wygenero-

wałaby w okresie pierwszym zwroty różne od zwrotu strategii X o pewną krańcową wielkość. Oznaczmy te zwroty poprzez

$$\frac{1}{B_0^0(1)} \left(\frac{2}{(k+1)} + \frac{2\varepsilon_1}{(k+1)} \right), \frac{1}{B_0^0(1)} \left(\frac{2}{(k+1)} - \frac{2\varepsilon_2}{(k+1)} \right), \frac{1}{B_0^0(1)} \left(\frac{2k}{(k+1)} + \frac{2k\varepsilon_3}{(k+1)} \right),$$

oraz $\frac{1}{B_0^0(1)} \left(\frac{2k}{(k+1)} - \frac{2k\varepsilon_4}{(k+1)} \right)$

gdzie $\varepsilon_1, \varepsilon_2, \varepsilon_3$ oraz ε_4 stanowią pewne dowolnie małe, niezerowe wielkości. W następnej kolejności wyznaczmy dla wszystkich możliwych ścieżek osiągalne przedziały zwrotów dla drugiego okresu. Ponieważ w każdym węźle drzewa ekstremalne poziomy zwrotów osiągalne są przy stuprocentowej alokacji majątku w obligacje o maksymalnym lub minimalnym terminie wykupu, osiągalne poziomy zwrotów w drugim okresie wynoszą odpowiednio:

- dla cen ze stanu 1 w pierwszym okresie i 2 w drugim okresie $\frac{B_0^0(1)}{B_0^0(2)} \left(l \frac{k+1}{2} + (1-l) \frac{2k}{k^2+1} \right)$,
- dla cen ze stanu 1 w pierwszym okresie i 1 w drugim okresie $\frac{B_0^0(1)}{B_0^0(2)} \left(l \frac{k+1}{2} + (1-l) \frac{2k^3}{k^2+1} \right)$,
- dla cen ze stanu 0 w pierwszym okresie i 1 w drugim okresie $\frac{B_0^0(1)}{B_0^0(2)} \left(l \frac{k+1}{2} + (1-l) \frac{2}{k^2+1} \right)$,
- natomiast dla cen ze stanu 0 zarówno w pierwszym, jak i drugim okresie $\frac{B_0^0(1)}{B_0^0(2)} \left(l \frac{k+1}{2} + (1-l) \frac{2k^2}{k^2+1} \right)$, gdzie l jest pewna liczba z przedziału $<0;1>$.

Spełnianie przez strategię X' wymogu zerowania wartości oczekiwanej z funkcji straty implikuje osiągalność takich zwrotów w drugim okresie, które przy danych zwrotach w okresie pierwszym satysfakcjonują, dla dowolnie małych $\varepsilon_1, \varepsilon_2, \varepsilon_3, \varepsilon_4$ oraz l z przedziału $<0;1>$, następujący układ nierówności

$$\left\{ \begin{array}{l} V_0 \frac{1}{B_0^0(2)} \left(\frac{2}{(k+1)} + \frac{2\varepsilon_1}{(k+1)} \right) \left(l \frac{k+1}{2} + (1-l) \frac{1+k}{k^2+1} \right) \geq L \\ V_0 \frac{1}{B_0^0(2)} \left(\frac{2}{(k+1)} - \frac{2\varepsilon_2}{(k+1)} \right) \left(l \frac{k+1}{2} + (1-l) \frac{k^3+k^3}{k^2+1} \right) \geq L \\ V_0 \frac{1}{B_0^0(2)} \left(\frac{2k}{(k+1)} + \frac{2k\varepsilon_3}{(k+1)} \right) \left(l \frac{k+1}{2k} + (1-l) \frac{1+k^{-1}}{k^2+1} \right) \geq L \\ V_0 \frac{1}{B_0^0(2)} \left(\frac{2k}{(k+1)} - \frac{2k\varepsilon_4}{(k+1)} \right) \left(l \frac{k+1}{2k} + (1-l) \frac{k+k^2}{k^2+1} \right) \geq L \end{array} \right. \quad (8)$$

Przy uwzględnieniu faktu, że zdyskontowana wartość zobowiązania równa się wartości początkowej majątku oraz przyjmując $\varepsilon = \min\{\varepsilon_1; \varepsilon_2; \varepsilon_3; \varepsilon_4\}$, po przekształceniach układ ten sprowadza się do

$$\left\{ \begin{array}{l} (1+\varepsilon) \left(l + (1-l) \frac{2}{(k^2+1)} \right) \geq 1 \\ (1-\varepsilon) \left(l + (1-l) \frac{2k^2}{(k^2+1)} \right) \geq 1 \end{array} \right. \quad (9)$$

Dla dowolnego ustalonego $l \in (0;1)$ wyrażenie w dużym nawiasie pierwszej nierówności jest mniejsze od jedności, a więc można dobrać odpowiednio małe ε , żeby nierówność nie zachodziła. Dla l równego 1 nie zachodzi druga nierówność.

Powyższa sprzeczność dowodzi nie wprost, że nie istnieje strategia inwestycyjna inna niż trywialna replikacja zwrotów z obligacji zapadającej w momencie płatności zobowiązania, która gwarantowałaby zerowanie się oczekiwanej wartości funkcji straty Cvitanica i Karatzasa. Przedstawione wyżej dowodzenie można łatwo uogólnić na dowolną skończoną liczbę okresów w modelu.

Interpretacja otrzymanego wyniku polega na spostrzeżeniu, że nieistnienie poszukiwanej strategii inwestycyjnej wynika intuicyjnie z nieparzystości funkcji straty Cvitanica i Karatzasa, która penalizując straty nie pozwala jednocześnie na kompensację oczekiwanych niedoborów majątku za pomocą oczekiwanych nadwyżek majątku. Przy uwzględnieniu specyfiki modelu Ho-Lee oznacza to, że każda strategia odbiegająca od replikacji zwrotu z obligacji zero-kuponowej zapadającej w horyzoncie zobowiązania (tzn. implikująca inny profil wypłat) prowadzi do niezerowej wartości funkcji straty Cvitanica i Karatzasa dla przynajmniej jednej ścieżki na drzewie ewolucji cen.

ANEKS

Aby dowieść prawdziwości wzoru (1.3), wykorzystujemy następujące równanie

$$B_t^n(m) = \prod_{j=t}^{m-1} g[j - (m-1)] B_j^n(j+1) \quad (\text{A.1})$$

Równanie (A.1) udowodnimy w sposób indukcyjny, wykazując jego prawdziwość dla $t=m-1$ i wykazując, że z jego prawdziwości dla $t+1$ wynika prawdziwość dla t . Podstawiając $t=m-1$ otrzymujemy

$$B_{m-1}^n(m) = g[(m-1) - (m-1)] B_{m-1}^n(m) = B_{m-1}^n(m)$$

co wykazuje prawdziwość pierwszego kroku indukcji. Korzystając z (1.2), zakładając prawdziwość (A.1) dla $t+1$ i pamiętając o (1.1) otrzymujemy

$$\begin{aligned} B_t^n(m) &= B_t^n(t+1) [q B_{t+1}^{n+1}(m) + (1-q) B_{t+1}^n(m)] = \\ &= B_t^n(t+1) \left[q \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^{n+1}(j+1) + (1-q) \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) \right] = \\ &= B_t^n(t+1) \left[q \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) \frac{1}{k} + (1-q) \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) \right] = \\ &= B_t^n(t+1) \left[\left(\frac{1}{k} \right)^{(m-1)-t} q \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) + (1-q) \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) \right] = \\ &= B_t^n(t+1) \left[q \left(\frac{1}{k} \right)^{(m-1)-t} + (1-q) \right] \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) = \\ &= B_t^n(t+1) g[t - (m-1)] \prod_{j=t+1}^{m-1} g[j - (m-1)] B_j^n(j+1) = \prod_{j=t}^{m-1} g[j - (m-1)] B_j^n(j+1) \end{aligned}$$

co kończy dowód indukcyjny wzoru (A.1). Aby wykazać, że z (A.1) wynika zależność (1.3), podstawiamy w nim $t=0$, $n=0$ oraz $t=0$, $m=t$, $n=0$ i dzielimy otrzymane wyrażenia przez siebie

$$\frac{B_0^0(m)}{B_0^0(t)} = \frac{\prod_{j=0}^{m-1} g[j - (m-1)] B_j^0(j+1)}{\prod_{j=0}^{t-1} g[j - (t-1)] B_j^0(j+1)} = \frac{B_t^0(m) \prod_{j=0}^{t-1} g[j - (m-1)] B_j^0(j+1)}{\prod_{j=0}^{t-1} g[j - (t-1)] B_j^0(j+1)} =$$

$$= B_t^0(m) \frac{\prod_{j=0}^{t-1} g[j-(m-1)]}{\prod_{j=0}^{t-1} g[j-(t-1)]}$$

co przy wykorzystaniu zależności (1.1) prowadzi bezpośrednio do (1.3).

Literatura

1. Cvitanic J., Karatzas I. (1999). On dynamic Measures of Risk. Finance & Stochastics.
2. Ho T., Lee S. (1986). Term Structure Movements and Pricing Interest Rate Contingent Claims, Journal of Finance.
3. Utkin J., Żułtak S. (2007). Miary niedotrzymania zobowiązania na rynku dyskretnym, referat na konferencję „Innowacje w finansach i ubezpieczeniach – metody matematyczne, ekonometryczne i informatyczne 2007”. AE, Katowice.

ANALIZA PROJEKTÓW

Paweł Błaszczyk

Maria B. Kania

Uniwersytet Śląski w Katowicach

Tomasz Błaszczyk

Akademia Ekonomiczna w Katowicach

DWUKRYTERIALNA ANALIZA PROJEKTU MODELOWANEGO Z UWZGLĘDNIENIEM BUFORÓW CZASOWYCH I KOSZTOWYCH

Wstęp

Problem analizy czasowo-kosztowej, pozwalającej na wyznaczenie takiego planu projektu, który satysfakcjonuje decydenta ze względu na potrzebę zaplanowania jak najwcześniejszego terminu zakończenia realizacji przy jak najmniejszym nakładzie środków jest jednym z podstawowych wielokryterialnych problemów planowania projektów. Pierwsze badania nad powyższym problemem, autorstwa Fulkersona [3] i Kelleya [7] opublikowane zostały w latach 60. XX wieku. Szczegółowy przegląd dotychczasowych wyników był wielokrotnie publikowany, między innymi w pracy innych [2]. Koncepcja podjęcia badań, których wyniki opisane są w niniejszej pracy wynika z rozważań nad możliwościami rozszerzenia na większą liczbę kryteriów podejścia opisanego przez Goldratta [4] i nazwanego łańcuchem krytycznym. Pierwotny opis metody nie wykorzystywał zapisu formalnego, jedynie werbalny. Metody kwantyfikacji łańcucha oraz buforów czasowych, których szczegółowy opis zawarty jest między innymi w pracy [10], są wkładem kolejnych autorów. Zagadnienia buforowania innych czas cech projektu omawiane były już między innymi w pracy [8; 5; 1]. Metoda łańcucha krytycznego, wielokrotnie będąca przedmiotem naukowej dyskusji [6; 9; 11] nie jest pozbawiona wad utrudniających jej praktyczne stosowanie, lecz poprzez uwzględnienie w nim czynnika behawioralnego ma pewną przewagę nad metodami stricte ilościowymi. Mianowicie, poprzez uwzględnienie wpływu człowieka na mierzalne cechy

projektu jesteśmy również w stanie wykorzystać go w celu poprawy tychże cech w zamian za określony ekwiwalent, najczęściej wyrażony w pieniądzu. Przykład takiego działania, wykorzystującego nadzwyczajny fundusz premiowy, opisany został w pracy [1]. Dalsza część niniejszej publikacji jest wynikiem kontynuacji badań nad próbami kwantyfikacji buforów kosztowych oraz przedstawienia problemu czasowo-kosztowego w ich kontekście. W odróżnieniu od dotychczasowych prac w przyjętej procedurze wyodrębnieniu przeszacowań podlegają podane przez podwykonawców poszczególnych czynności nakłady pracy wymagane do ich pełnej realizacji. Zaproponowany w pracy model zakłada możliwość zachęcenia wykonawców prac w projekcie do partycypacji w ryzyku niezrealizowania mniej ostrożnych szacunków w zamian za udział w korzyściach wynikających z ewentualnej szybszej i tańszej realizacji. Koncepcja teoretyczna zilustrowana jest przykładem liczbowym, opracowanym z wykorzystaniem arkusza kalkulacyjnego.

1. Model teoretyczny

Rozważmy projekt składający się z n czynności $x_1, \dots, x_n$. Każdą z czynności x_i charakteryzuje zarówno czas, jak i koszt jej wykonania. Założmy, że zarówno na czas, jak i na koszt wykonania całego projektu wpływa co najwyżej q czynników. Rozważmy także macierz X określoną następująco

$$X = \begin{bmatrix} x_{11} & \cdots & x_{1q} \\ \vdots & \ddots & \vdots \\ x_{n1} & \cdots & x_{nq} \end{bmatrix}$$

Elementy macierzy X przyjmują wartości 0 lub 1, przy czym, jeżeli element x_{ij} macierzy X ma wartość 1, to mówimy o wpływie czynnika j na czynność x_i . W przeciwnym wypadku czynnik j nie wpływa na czynność x_i . Macierz X nazywamy macierzą czynników.

Niech $K = [k_{ij}]_{i=1, \dots, n; j=1, \dots, q}$ będzie macierzą zawierającą współczynniki jednostkowego kosztu q czynników dla $x_1, \dots, x_n$ czynności. Ponadto, niech $W^m = [w_1^m, \dots, w_n^m]$ będzie wektorem minimalnego wymaganego nakładu pracy dla czynności $x_1, \dots, x_n$. Na podstawie macierzy czynników X oraz wektora W^m obliczamy całkowity nakład pracy w_i dla czynności x_i zgodnie ze wzorem

$$w_i = f_{w_i} (x_{i1}, \dots, x_{iq}, w_i^m)$$

gdzie f_{w_i} jest pewną funkcją. Co więcej założmy, że dany jest wektor $R = [r_1, \dots, r_q]$ zawierający ograniczenie dostępności poszczególnych czynników. Niech $T = [t_{ij}]_{i=1, \dots, n; j=1, \dots, q}$ będzie macierzą nakładu pracy poszczególnych czynników w każdej czynności. Na podstawie macierzy X, T oraz K obliczamy koszt i czas trwania każdej czynności zgodnie ze wzorami

$$k_i = f_{i_k} (x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}, k_{i1}, \dots, k_{iq})$$

oraz

$$t_i = f_{i_t} (x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq})$$

gdzie f_{i_k}, f_{i_t} są pewnymi funkcjami. Powyższe funkcje będziemy nazywać odpowiednio funkcją kosztu oraz funkcją czasu. Wobec tego całkowity koszt oraz całkowity czas trwania projektu są dane odpowiednio wzorami

$$K_c = \sum_{i=1}^n k_i$$

oraz

$$T_c = \max_{i=1, \dots, q} (ES_i + t_i)$$

gdzie ES_i jest najwcześniejszym czasem rozpoczęcia czynności x_i . Pod powyższymi warunkami minimalizujemy całkowity koszt projektu. Sprowadza się to do znalezienia odpowiednich nakładów pracy wszystkich czynników w poszczególnych czynnościach. Ze zbioru rozwiązań optymalnych wybieramy te, dla którego czas trwania projektu jest najmniejszy. W ten sposób otrzymujemy rozwiązanie optymalne dla oszacowań bezpiecznych. Ze względów bezpieczeństwa realizacji projektu nakłady pracy poszczególnych czynności mogą być przeszacowane, co w konsekwencji prowadzi do przeszacowania kosztu i czasu realizacji projektu. Stąd otrzymujemy

$$k_i = k_i^e + k_i^B$$

$$t_i = t_i^e + t_i^B$$

gdzie k_i^e, t_i^e są odpowiednio realnym kosztem i realnym czasem realizacji czynności x_i , natomiast k_i^B, t_i^B są odpowiednio ukrytym zapasem kosztu oraz ukrytym zapasem czasu realizacji czynności x_i . Stąd możemy zapisać, że całkowity koszt oraz całkowity czas realizacji projektu dane są wzorami

$$K_c = K^e + K^B$$

oraz

$$T_c = T^e + T^B$$

gdzie K^e, T^e są odpowiednio rzeczywistym kosztem i czasem realizacji projektu, natomiast K^B, T^B są odpowiednio buforem kosztu oraz buforem czasu realizacji projektu. W celu wyznaczenia buforów K^B, T^B musimy obliczyć nakład pracy, który zdaniem decydenta jest uzasadniony dla poszczególnych czynności. Obliczamy go dokonując zmian wartości odpowiednich elementów w macierzy X oraz wykorzystując funkcję w_i dla każdej czynności x_i . W ten sposób otrzymujemy macierz czynników X^* oraz wektor nakładów pracy W^* . Następnie stosujemy omówioną wyżej procedurę, jednak pod dodatkowymi warunkami

$$t_{ij} \geq t_{ij}^* \text{ dla } i = 1, \dots, n; j = 1, \dots, q$$

gdzie $T^* = [t_{ij}^*]$ jest macierzą nakładów pracy dla każdego czynnika w poszczególnych czynnościach, obliczoną dla macierzy X^* . Ponieważ jednak wystąpienie wszystkich czynników jednocześnie jest mało prawdopodobne, możemy zredukować bufor projektu. Wówczas

$$K_r^B = \alpha K^B$$

$$T_r^B = \beta T^B$$

gdzie $\alpha, \beta \in \langle 0, 1 \rangle$ są współczynnikami redukującymi wielkość bufora.

$$K^P = K^e + K_r^B$$

$$T^P = T^e + T_r^B$$

Część zaoszczędzonych pieniędzy może być przeznaczona na tzw. fundusz premiowy B i być rozdzielona pomiędzy poszczególne czynniki. W tym celu wprowadzimy wektor wag $S = [s_i]_{i=1, \dots, n}$ określający istotność poszczególnych czynności w projekcie, przy czym $s_i \in [0, 1]$. W celu podziału fun-

duszu premiowego definiujemy nową funkcję f_{i_b} zależną od wartości zaoszczędzonego nakładu pracy, ważności czynności, oraz od tego czy czynność x_i jest czynnością krytyczną czy też nie. W ogólnym przypadku czynnik i dla $i \in \{1, \dots, q\}$ może otrzymać premię b_i zgodnie z poniższym wzorem

$$b_i = f_{i_b}(s_i, D_i^W, c, D_i^K, D_i^T)$$

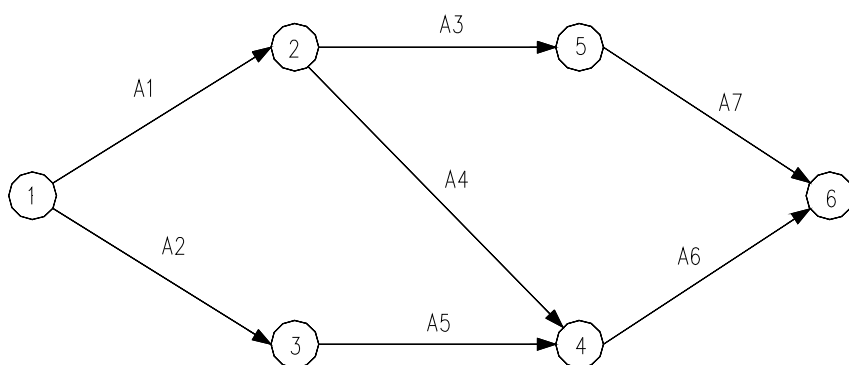
gdzie s_i jest wagą, D_i^W jest wielkością zaoszczędzonego nakładu pracy, D_i^K jest wartością zaoszczędzonego nakładu kosztu, D_i^T jest wartością zaoszczędzonego nakładu czasu dla czynności x_i , zmienna c przyjmuje wartości 0 lub 1 w zależności od tego, czy czynność x_i znajduje się na ścieżce krytycznej czy też nie. Natomiast f_{i_b} jest pewną funkcją. W omawianym przypadku wykorzystaliśmy następującą funkcję

$$b_i = \begin{cases} \frac{s_j D_j^W}{s^1 D^1} \gamma_1 B & \text{dla } x_i \text{ leżących na ścieżce krytycznej} \\ \frac{s_j D_j^W}{s^2 D^2} \gamma_2 B & \text{w przeciwnym przypadku} \end{cases}$$

gdzie B jest funduszem premiowym, $\gamma_2 < \gamma_1$, $\gamma_1 + \gamma_2 = 1$, s^1 jest sumą wag istotności czynności leżących na ścieżce krytycznej, zaś s^2 jest odpowiednią sumą dla czynności nie leżących na ścieżce krytycznej, D_j^W jest sumą zaoszczędzonych nakładów pracy dla czynności x_i , D^1 jest sumą zaoszczędzonych nakładów pracy dla czynności leżących na ścieżce krytycznej, natomiast D^2 jest sumą zaoszczędzonych nakładów pracy dla czynności nie leżących na ścieżce krytycznej.

2. Przykład obliczeniowy

Realizacja zadania wymaga wykonania siedmiu czynności $A_1, A_2, \dots, A_7$. Sieć zależności pomiędzy czynnościami przedstawia rys. 1.



Rys. 1. Graf projektu analizowanego w przykładzie obliczeniowym.

Nakład pracy W niezbędny do realizacji poszczególnych czynności przedstawia tabela 1.

Tabela 1

Nakład pracy niezbędny do realizacji poszczególnych czynności

Czynność	Nakład pracy
A_1	200
A_2	240
A_3	150
A_4	400
A_5	340
A_6	200
A_7	690

Suma nakładu pracy potrzebnej do realizacji projektu wynosi 2200 roboczogodzin. Do realizacji projektu przeznaczono 4 pracowników. Ich możliwy udział w poszczególnych czynnościach został zapisany za pomocą macierzy czynników X i przedstawia go tabela 2.

Tabela 2

Możliwy udział pracowników w poszczególnych czynności

Czynność	Pracownik 1	Pracownik 2	Pracownik 3	Pracownik 4	Pracownik 5
A ₁	0	0	1	1	0
A ₂	1	1	0	0	0
A ₃	0	0	0	0	1
A ₄	0	0	1	1	0
A ₅	1	1	0	0	0
A ₆	0	1	0	1	0
A ₇	1	0	1	0	1

Natomiast macierz kosztów K została przedstawiona za pomocą tabeli 3.

Tabela 3

Macierz kosztów

Czynność	Pracownik 1	Pracownik 2	Pracownik 3	Pracownik 4	Pracownik 5
A ₁	50	60	33	35	28
A ₂	45	50	52	48	35
A ₃	32	40	30	68	43
A ₄	67	55	63	57	47
A ₅	55	60	45	54	51
A ₆	45	42	50	50	38
A ₇	50	50	45	57	49

Ograniczenia nakładu pracy R dla poszczególnych pracowników przedstawione są za pomocą tabeli 4.

Tabela 4

Ograniczenia nakładu pracy dla poszczególnych pracowników

Pracownik 1	Pracownik 2	Pracownik 3	Pracownik 4	Pracownik 5
600	450	650	200	350

W celu usprawnienia przebiegu obliczeń do rozwiązania zadania wykorzystano arkusz kalkulacyjny Microsoft EXCEL. Skonstruowanie modelu projektu pozwoliło na wyznaczenie macierzy nakładu pracy T prezentowanej w tabeli 5.

Tabela 5

Macierz nakładu pracy dla poszczególnych pracowników i czynności

Czynność	Pracownik 1	Pracownik 2	Pracownik 3	Pracownik 4	Pracownik 5
A ₁	0	0	200	0	0
A ₂	118	122	0	0	0
A ₃	0	0	0	0	150
A ₄	0	0	200	200	0
A ₅	242	98	0	0	0
A ₆	0	200	0	0	0
A ₇	240	0	250	0	200

Koszt k_i każdej z czynności został oszacowany z wykorzystaniem następującej funkcji

$$k_i = f_{i_k} (x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}, k_{i1}, \dots, k_{iq}) = \sum_{j=1}^q x_{ij} k_{ij} t_{ij}$$

Czas t_i trwania każdej czynności został oszacowany z wykorzystaniem funkcji

$$t_i = f_{i_t} (x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}) = \max_{1 \leq j \leq q} \{x_{iq} t_{iq}\}$$

Przy powyższym bezpiecznym podziale pracy pomiędzy dostępnych pracowników całkowity koszt projektu K_c został oszacowany na $K_c = 109100$ jednostek pieniężnych (j.p.), natomiast całkowity czas realizacji projektu T_c został oszacowany na $T_c = 600$ godzin. Powyższe wyniki uzyskano dla bezpiecznych oszacowań nakładów pracy. W wyniku analizy nakładów pracy niezbędnych do realizacji poszczególnych czynności potrzymano nowe, realne, wartości oszacowań nakładów pracy W^* , które przedstawiono w tabeli 6.

Tabela 6

Wektor nakładów pracy dla poszczególnych czynności

Czynność	Nakład
A ₁	180
A ₂	216
A ₃	135
A ₄	360
A ₅	306
A ₆	180
A ₇	601

Ponowne wykorzystanie arkusza kalkulacyjnego Microsoft EXCEL pozwoliło na uzyskanie nowego podziału nakładów pracy dla poszczególnych pracowników i czynności. Otrzymana nowa macierz nakładu pracy T^* została przedstawiona na tabeli 7.

Tabela 7

Macierz realnych nakładu pracy dla poszczególnych pracowników i czynności

Czynność	Pracownik 1	Pracownik 2	Pracownik 3	Pracownik 4	Pracownik 5
A ₁	0	0	180	0	0
A ₂	118	98	0	0	0
A ₃	0	0	0	0	135
A ₄	0	0	160	200	0
A ₅	242	64	0	0	0
A ₆	0	180	0	0	0
A ₇	151	0	250	0	200

Przy powyższym podziale pracy pomiędzy pracowników całkowity koszt projektu K_e został oszacowany na $K_e = 96745$, natomiast całkowity czas realizacji projektu T_e został oszacowany na $T_e = 565$ godzin. Oszczędności w stosunku do wariantu bezpiecznego wynoszą odpowiednio $K^B = 12355$ j.p. dla kosztu i $T^B = 35$ godzin dla czasu. Stosując współczynniki zmniejszające wielkość buforów $\alpha = 0,5$ dla kosztu oraz $\beta = 0,6$ dla czasu otrzymujemy zredukowane buforu kosztu $K_r^B = 6177,5$ j.p. oraz czasu godzin $T_r^B = 21$. Wówczas koszt i czas realizacji projektu wynoszą odpowiednio $K^P = 102922,5$ oraz $T^P = 586$. Oszacowane wartości kosztu i czasu powinny być wystarczające do prawidłowego zrealizowania projektu. Część z pozostałej różnicy pomiędzy kosztem projektu K^P a kosztem w wariantcie bezpiecznym K_e może tworzyć tzw. fundusz premiiowy B . Dla podziału funduszu premiiowego pomiędzy poszczególnych pracowników przyjęto wektor wag S przedstawiony w tabeli 8.

Tabela 8

Wektor wag dla poszczególnych czynności

Czynność	waga
A ₁	0,7
A ₂	0,3
A ₃	0,4
A ₄	0,7
A ₅	0,4
A ₆	0,7
A ₇	0,5

Ścieżką krytyczną jest ścieżka A_1, A_4, A_6 o czasie przejścia 630 godzin. Wartości zaoszczędzonych nakładów pracy D_j^W , czasu D_j^T i kosztów D_j^K przedstawia tabela 9.

Tabela 9

Wartości zaoszczędzonych nakładów pracy, czasu i kosztu dla poszczególnych czynności

Czynność	Nakład pracy	Nakład kosztu	Nakład czasu
A_1	20	660	20
A_2	24	1200	4
A_3	15	645	35
A_4	40	2520	20
A_5	34	2040	4
A_6	20	840	40
A_7	89	4450	35
Suma	242	12355	158

Wykorzystana funkcja podziału premii jest postaci

$$b_i = \begin{cases} \frac{s_j D_j^W}{s^1 D^1} \gamma_1 B & \text{dla } x_i \text{ leżących na ścieżce krytycznej} \\ \frac{s_j D_j^W}{s^2 D^2} \gamma_2 B & \text{w przeciwnym przypadku} \end{cases}$$

przy czym przyjęto współczynniki $\gamma_1 = 0,8$ oraz $\gamma_2 = 0,2$, natomiast D^1, D^2, s^1, s^2 są odpowiednio równe $D^1 = 80$, $D^2 = 162$, $s^1 = 2,1$, $s^2 = 1,6$. Przy powyższych założeniach otrzymujemy następujący podział funduszu premialnego zawarty w tabeli 10.

Tabela 10

Wektor wag dla poszczególnych czynności

Czynność	Udział w funduszu premialnym
A_1	411,83
A_2	34,32
A_3	28,60
A_4	823,67
A_5	64,83
A_6	411,83
A_7	212,11

Podział premii pomiędzy poszczególnych pracowników zaangażowanych w projekt odbywa się proporcjonalnie do podziału puli przypadającej na daną czynność pomiędzy pracowników odpowiedzialnych za przyspieszenie danej czynności

Podsumowanie

W zagadnieniach planowanie projektów istnieje możliwość wyodrębniania zapasów bezpieczeństwa ukrytych w szacunkach, które mogą dotyczyć nie tylko spodziewanego czasu realizacji. Wyniki wcześniejszych badań wskazują, że mechanizm ten znajduje zastosowanie również w procesach budżetowania projektów. Przyjęta w niniejszej pracy teza stanowiąca o istnieniu przeszacowań w prognozach nakładów pracy wydaje się również uzasadniona. Przeprowadzone na gruncie teoretycznym rozważanie jest zgodne z omówionym w pracy [1] podejściem dotyczącym buforowania kosztu projektu co do procedury wyznaczania buforów oraz podziału osiągniętych korzyści. Podejście to zostało jednakże rozszerzone o kwestię czynników, które mogą mieć lub nie mieć wpływu na realizację poszczególnych czynności. Hipotetyczna możliwość osiągnięcia wymiernych korzyści z zastosowania proponowanej metody wykazana została w przykładzie obliczeniowym. Podkreślić jednakże należy, że udowodnienie sprawności działania przedstawionej koncepcji w warunkach rzeczywistych wymaga praktycznego jej zastosowania w rzeczywistym projekcie. Przeprowadzenie weryfikacji praktycznej będzie celem dalszych badań nad podjętym zagadnieniem.

Literatura

1. Błaszczyk T., Nowak B. (2008). Szacowanie kosztów projektu w oparciu o metodę łańcucha krytycznego. W: Modelowanie preferencji a ryzyko '08. Red. T. Trzaskalik. AE, Katowice.
2. Brucker P., Drexel A., Möhring R., Neumann K., Pesch E. (1999). Resource-constrained Project Scheduling: Notation, Classification, Models, and Methods. *European Journal of Operational Research* 112, 3-41.
3. Fulkerson D.R. (1961). A Network Flow Computation for Project Cost Curves. *Management Science*, 7, 167-178.
4. Goldratt E. (1997). *Critical Chain*. North River Press.
5. Gonzalez V., Alarcon L.F., Molenaar K. (2009). Multiobjective Design of Work-In-Process Buffer for Scheduling Repetitive Projects. *Automation in Construction*, 18, 95-108.

6. Herroelen W., Leus R. (2001). On the Merits and Pitfalls of Critical Chain Scheduling. *Journal of Operations Management*, 19, 559-577.
7. Kelley J.E. (1961). Critical-path Planning and Scheduling: Mathematical Basis. *Operations Research*, 9, 296-320.
8. Leach L. (2003). Schedule and Cost Buffer Sizing: How Account for the Bias between Project Performance and Your Model. *Project Management Journal*, 34 (2), 34-47.
9. Rogalska M., Bożejko W., Hejducki Z. (2008). Time/cost Optimization Using Hybrid Evolutionary Algorithm in Construction Project Scheduling. *Automation in Construction*, 18, 24-31.
10. Tukul O.I., Rom W.O., Eksioglu S.D. (2006). An Investigation of Buffer Sizing Techniques in Critical Chain Scheduling. *European Journal of Operational Research*, 172, 401-416.
11. Van de Vonder S., Demeulemeester E., Herroelen W., Leus R. (2005). The Use of Buffers in Project Management: The Trade-off Between Stability and Makespan. *International Journal of Production Economics*, 97, 227-240.

Monika Fedyczak

Dorota Kuchta

Politechnika Wrocławska

ZRÓWNOWAŻONE PODEJŚCIE W ZARZĄDZANIU RYZYKIEM PROJEKTÓW EUROPEJSKICH

Wstęp

Fundusze europejskie to temat szeroko dyskutowany w mediach i ważny dla społeczeństwa. Są od pewnego czasu dostępne dla naszego kraju w dość znacznych ilościach, jednak problemem jest ich uzyskanie, a potem właściwe wykorzystanie. Wiąże się to z koniecznością właściwego zarządzania ogromnymi projektami zarówno pod względem zakresu, jak i kosztów. Każdy, kto kiedykolwiek zarządzał projektem wie, że jest to pewnego rodzaju sztuka, wymagająca znacznych umiejętności, doświadczenia i że nawet umiejętności i doświadczenie nie zapewniają pełnej ochrony przed niepowodzeniem. Projekty są bowiem z samej swojej natury związane z niepewnością i ryzykiem. Dotyczy to również projektów współfinansowanych przez Unię (tzw. projektów europejskich), a osoby odpowiedzialne za ich realizację nie zawsze mają wystarczającą wiedzę i doświadczenie z zakresu zarządzania projektami. Ponadto, każdy typ projektu wymaga innego podejścia, innych metod zarządzania, innego rodzaju doświadczenia. Dla projektów europejskich nie ma jeszcze wypracowanych procedur i metod, jak dla np. projektów informatycznych, choć koszty niepowodzenia projektu współfinansowanego przez Unię mogą być z punktu widzenia społeczeństwa nieporównywalnie wyższe niż koszty niepowodzenia jednego z projektów realizowanych przez firmę informatyczną. Dlatego problem zarządzania projektami europejskimi wydaje się ważnym tematem badawczym.

W niniejszym opracowaniu zaprezentujemy koncepcję metody zrównoważonego (tzn. rozpatrywanego z różnych perspektyw) zarządzania ryzykiem projektów współfinansowanych przez Unię Europejską. Przedtem podamy definicję i przedstawimy specyfikę projektu europejskiego oraz przypomnimy podstawowe informacje dotyczące ogólnie zarządzania ryzykiem projektowym.

Omówimy również przykład zrealizowanego już projektu europejskiego, który podkreśli konieczność opracowania metody wspomagającej zrównoważone zarządzanie projektami europejskimi.

1. Podstawowe informacje dotyczące projektów europejskich

Definicja projektu europejskiego mówi, iż jest to przedsięwzięcie, które ma służyć realizacji celu określonego w programie, będące przedmiotem umowy o dofinansowanie między beneficjentem a instytucją zarządzającą, instytucją wdrażającą albo działającą w imieniu instytucji zarządzającej instytucją pośredniczącą. Projekt europejski jest to najmniejsza dająca się wyodrębnić jednostka stanowiąca przedmiot pomocy. Program natomiast to dokument realizowany w ramach polityki strukturalnej państwa, przyjęty przez Radę Ministrów i Komisję Europejską, składający się z zestawienia priorytetów oraz wieloletnich działań, które mogą być wdrażane poprzez jeden lub kilka Funduszy Europejskich, jeden lub kilka innych dostępnych instrumentów finansowych oraz Europejski Bank Inwestycyjny [1].

Cechą odróżniającą projekty europejskie od innych rodzajów projektów jest to, iż planowanie oraz realizacja takich projektów wymaga stosowania sformalizowanych reguł i regulacji europejskich, w związku z tym dowolność planowania i realizacji takich projektów jest ograniczona [6, s. 14]. Charakterystyczną cechą projektów europejskich jest również to, iż w ich realizacji uczestniczy wiele podmiotów (projektodawcy, decydenci, wykonawcy oraz beneficjenci), które zazwyczaj nie mają wiedzy i doświadczenia związanego z zarządzaniem projektami, a w szczególności z zarządzaniem ryzykiem projektowym. Uczestnicy projektu to podmioty z różnych środowisk, reprezentujące różne poziomy i obszary wiedzy, o zróżnicowanym doświadczeniu w realizacji projektów [6, s. 14]. Po 2004 roku w Polsce projekty europejskie realizowane są w dużej liczbie i angażują dużą ilość środków (łączna suma środków w latach 2007-2013 wyniesie około 85,6 mld euro, z czego 67,3 mld euro będzie pochodziło z budżetu UE).

Projekty europejskie to przede wszystkim projekty budowlane, projekty dotyczące budowy i modernizacji obiektów inżynierii lądowej lub wodnej, obiektów użyteczności publicznej, szpitali, obiektów sportowych, rekreacyjnych oraz wypoczynkowych, budynków szkolnych, budynków szkół wyższych oraz budynków administracji publicznej. Celem projektów europejskich jest zwiększenie atrakcyjności regionu, ułatwienie funkcjonowania obywatelom. Publiczne projekty europejskie to projekty, które współfinansowane są przez obywateli i właśnie obywatelom ma przynieść wymierne korzyści ich realizacja.

Nie mają one charakteru komercyjnego – nie są nastawione na zysk, można określić je jako projekty non-profit, realizujące cele pozadochodowe. Realizacja tego typu projektów wymaga więc pełnej akceptacji projektu przez mieszkańców i lokalnych przedsiębiorców oraz determinacji władz miasta lub regionu do jego realizacji.

Realizacja projektów europejskich wiąże się z większą presją i kontrolą społeczną niż w przypadku projektów sektora prywatnego. Wynika to z faktu, iż realizacja tych przedsięwzięć przebiega w sposób bardziej jawny, a dodatkowo cel projektu oraz korzyści związane z jego realizacją dotyczą bezpośrednio społeczeństwa. Ponadto, projekty te poddane są bardzo ścisłej kontroli komisji europejskiej. W związku z tym niezwykle ważne jest zakończenie projektu sukcesem rozumianym, nawiasem mówiąc, często w różny sposób przez lokalną społeczność, komisję europejską i pozostałych udziałowców projektu. Dlatego zarządzanie ryzykiem projektowym europejskich projektów publicznych jest niezwykle istotne.

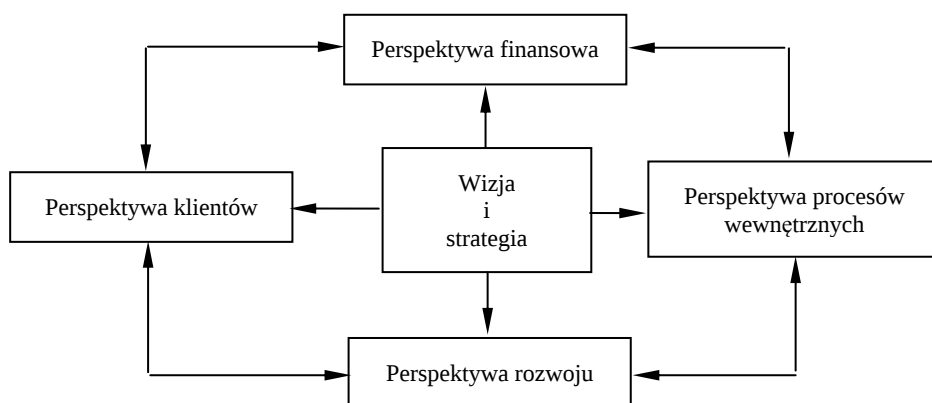
Przeprowadzone wywiady z kierownikami projektów europejskich wskazują podstawowe aspekty i problemy, na które trzeba zwrócić uwagę podczas realizacji projektów europejskich. Można tu wymienić między innymi:

- ochronę środowiska i zdrowia ludzi – hałas, powietrze, zieleń,
- ochronę dziedzictwa kulturowego i zabytków kultury współczesnej,
- ochronę interesu osób trzecich – konieczność zapewnienia dostępu do dróg, wody, kanalizacji, środków łączności, minimalizowanie hałasu, wibracji, zakłóceń elektrycznych, promieniowania,
- zapewnienie zgodności z normami technicznymi oraz dokumentacją techniczną
- konieczność zgromadzenia kompletnej dokumentacji – decyzja o warunkach zabudowy, decyzja o zezwoleniu na budowę nowego mostu, jezdni itp., decyzja o zwolnieniu od zakazu budowy w obrębie wałów przeciwpowodziowych,
- utrzymanie budowli i urządzeń we właściwym stanie,
- zaspokojenie ewentualnych roszczeń osób trzecich za wyrządzone w trakcie realizacji projektu szkody,
- konieczność przywrócenia terenu do stanu pierwotnego,
- konieczność uwzględnienia skomplikowanych zasad własności terenu – działki i tereny, na których realizowany jest projekt mogą być własnością wielu podmiotów, zarówno publicznych (Skarb Państwa), gmina, jak i prywatnych,
- organizację ruchu zastępczego i docelowego,
- konieczność dostosowania architektonicznego, zachowania zabytkowego charakteru obiektów,
- zachowanie ostrożności przy pracach – ochrona znaków geodezyjnych,

- uzgodnienia z właścicielami światłowodów, linii energetycznych, telekomunikacyjnych,
- zabezpieczenie środków niezbędnych do zrealizowania projektu.

2. Podstawowe informacje dotyczące Zrównoważonej Karty Wyników

Zrównoważona Karta Wyników to opracowana przez Kaplana i Nortona [2, s. 29] metoda kompleksowego pomiaru efektów działalności przedsiębiorstw, wychodząca z założenia, że przy pomiarze efektów działalności przedsiębiorstwa trzeba wykroczyć poza tradycyjne wskaźniki finansowe dotyczące przeszłości, że trzeba brać pod uwagę różne aspekty, aspekt finansowy i niefinansowy, przeszłość i przyszłość, punkt widzenia zarządu przedsiębiorstwa i jej klientów itp. Nawet jeśli w momencie zaproponowania Zrównoważonej Karty Wyników wskaźniki niefinansowe były w wielu przedsiębiorstwach uwzględniane, to było to robione w sposób nieuporządkowany, nieusystematyzowany, przypadkowy. Istotą Zrównoważonej Karty Wyników jest zrównoważone traktowanie wspomnianych aspektów i punktów widzenia. Metoda Zrównoważonej Karty Wyników opiera się na czterech perspektywach (rys. 1), które są ze sobą wzajemnie powiązane, a ich analiza powinna przebiegać w sposób zrównoważony tak, aby cele poszczególnych perspektyw tworzyły logiczny łańcuch przyczynowo-skutkowy i żeby żadna z perspektyw nie została zaniedbana kosztem innej.



Rys. 1. Zrównoważona Karta Wyników

Źródło: [2].

Każda z czterech perspektyw (finansowa, klienta, procesów wewnętrznych oraz rozwoju) charakteryzuje się taką samą strukturą. W każdej perspektywie znajduje się zestaw celów, które trzeba zrealizować, aby osiągnąć sukces. Cele tworzą łańcuch przyczynowo-skutkowy i są określone w taki sposób, aby uwzględniały wszystkie aspekty działalności organizacji. Do każdego celu stworzony jest zestaw mierników, które mają mierzyć sposób realizacji celu.

Przykładowe cele i przyporządkowane im mierniki to [2, s. 32, 75, 113, 115, 118, 124]:

- w perspektywie finansowej: wzrost przychodów – mierniki: stopa wzrostu sprzedaży w poszczególnych segmentach, udział przychodów z nowych produktów i usług,
- w perspektywie klienta: utrzymanie klientów – mierniki: lojalność klientów (roczny przyrost zakupów), stopień utrzymania trwałych relacji z klientami,
- w perspektywie procesów wewnętrznych: poprawa jakości procesów wewnętrznych – mierniki: liczba wadliwych sztuk na milion produktów, liczba zwrotów, liczba dobrych wyrobów do ogólnej liczby wyprodukowanych wyrobów,
- w perspektywie rozwoju: poprawa potencjału kadrowego przedsiębiorstwa – mierniki: miernik rotacji pracowników (odsetek pracowników, którzy odeszli z kluczowych stanowisk), miernik wydajności pracowników (przychód na jednego zatrudnionego, przychody ze sprzedaży do sumy wynagrodzeń), mierniki satysfakcji pracowników (ogólne zadowolenie z pracy w przedsiębiorstwie).

Wszystkie cele i mierniki powinny być tak skonstruowane, aby były ze sobą powiązane i tworzyły łańcuch przyczynowo-skutkowy, aby ułatwić koordynację działań i realizację strategii.

3. Podstawowe informacje dotyczące zarządzania ryzykiem projektu

Zarządzanie ryzykiem projektu to część procesu zarządzania projektem, w ramach którego realizowane są procesy identyfikacji, analizy i opracowania metod reagowania na ryzyko. Proces zarządzania ryzykiem składa się z sześciu faz [5, s. 4]:

1. Planowania procesu zarządzania ryzykiem – opracowania planu zarządzania ryzykiem dla konkretnego projektu.

2. Identyfikacji ryzyka – opisu zdarzeń, które mogą mieć negatywny na realizowany projekt, opisu źródeł ryzyka.

3. Kwalifikacji ryzyka – oceny ryzyka za pomocą metod jakościowych.

4. Pomiaru (kwantyfikacji) ryzyka – analizy ryzyka za pomocą metod ilościowych, pomiaru prawdopodobieństwa oraz skutków najważniejszych rodzajów ryzyka.

5. Planowania sposobów reakcji na ryzyko – oceny i komunikacji strategii przeciwdziałania, zapobiegania lub minimalizowania ryzyka.

6. Monitorowania i kontroli ryzyka – wdrażania metod zarządzania ryzykiem i planowania sposobów reagowania na ryzyko.

Aby zarządzanie ryzykiem było procesem efektywnym, musi być przeprowadzane systematycznie, a dodatkowo konieczne jest stosowanie sformalizowanych metod. Dla każdej z faz literatura podaje odpowiednie metody, są one jednak zazwyczaj bardzo ogólne (niedostosowane do specyfiki projektów europejskich) lub specyficzne dla innego rodzaju projektów (np. projektów informatycznych). Jeśli chodzi o zarządzanie ryzykiem projektów europejskich brak jest sformalizowanych metod wspomagających zarządzanie ryzykiem projektowym.

4. Aktualny stan zarządzania ryzykiem projektów europejskich

Wywiady przeprowadzone z koordynatorami projektów europejskich w Urzędzie Miejskim Wrocławia oraz obserwacje realizacji projektów europejskich pozwalają stwierdzić, iż proces zarządzania ryzykiem w projektach europejskich jest procesem niezbędnym do właściwej, przemyślanej realizacji projektu. Jednakże w projektach europejskich proces zarządzania ryzykiem jest traktowany marginalnie. Koordynatorzy projektów przyznają, iż osoby zajmujące się koordynacją projektów europejskich nie posiadają umiejętności zarządczych, wymaganych do realizacji tego typu przedsięwzięć. Przede wszystkim w projektach europejskich, gdzie kierownikami projektu są urzędnicy biura zarządzania projektami europejskimi, koordynatorzy są zaangażowani w realizację zbyt wielu projektów, a przez to nadmiernie przeciążeni, a także nie mają umiejętności, wiedzy i doświadczenia niezbędnych do właściwej identyfikacji, analizy oraz przeciwdziałania ryzyku projektowemu.

Dodatkowo, w projektach europejskich nie są stosowane sformalizowane metody zarządzania ryzykiem. Skutkuje to sytuacją, w której żadna faza zarządzania ryzykiem – planowanie zarządzania ryzykiem, identyfikacja ryzyka, jakościowa i ilościowa analiza ryzyka, planowanie reakcji na ryzyko oraz monitorowanie i kontrolowanie ryzyka – nie jest prowadzona w sposób uporządkowany i usystematyzowany, a w rzeczywistości niektóre (jeśli nie wszystkie) nie są przeprowadzane w ogóle. W tej sytuacji project manager nie jest w stanie skutecznie i efektywnie identyfikować ryzyka projektowego, źródeł ryzyka, nie

ma wiedzy o wpływie poszczególnych zdarzeń na wynik końcowy projektu, nie jest w stanie zaplanować strategii reakcji na ryzyko ani określić skutków wystąpienia poszczególnych zagrożeń.

Działania prowadzone przez project managerów nie są więc usystematyzowane, a ryzyko (jeśli jest analizowane) rozpatrywane jest wybiórczo. Selektywna analiza ryzyka powoduje, iż kierownik projektu nie widzi powiązania z innymi płaszczyznami i aspektami projektu, nie stara się równoważyć ryzyka, a tylko reaguje na pojawiające się na bieżąco zagrożenia.

W zarządzaniu projektami europejskimi istnieje konieczność zarządzania ryzykiem projektowym, jednakże aby działania te były skuteczne, konieczne jest znalezienie rozwiązań, które w sposób przejrzysty i łatwy do wdrożenia pozwolą kierownikom projektów europejskich na efektywniejszą identyfikację, analizę i przeciwdziałanie ryzyku projektowemu w celu zmniejszenia prawdopodobieństwa porażki projektu.

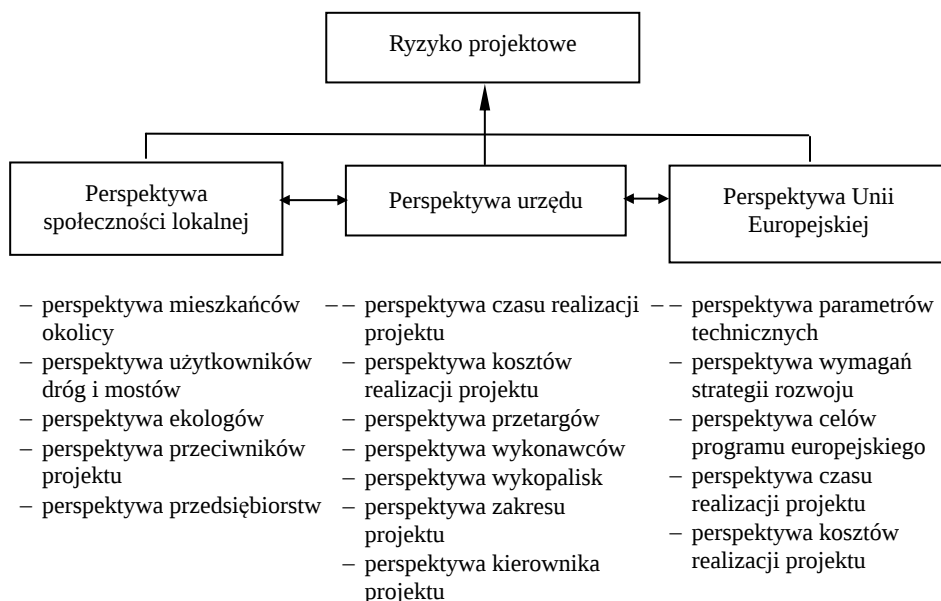
5. Zrównoważone podejście do zarządzania ryzykiem projektów europejskich

Propozycja zrównoważonego podejścia do zarządzania ryzykiem projektowym wykorzystuje niektóre elementy Zrównoważonej Karty Wyników. Propozycja opiera się głównie na opracowaniu perspektyw charakterystycznych dla projektów europejskich, a także opracowaniu zależności pomiędzy perspektywami oraz relacji pomiędzy źródłami ryzyka w poszczególnych perspektywach. W każdej perspektywie opracowane zostaną źródła ryzyka, cele zapewniające uniknięcie ryzyka bądź zmniejszenie jego skutków oraz mierniki stopnia realizacji celów, w taki sposób, aby analiza i pomiar ryzyka przeprowadzona była w sposób zrównoważony. Pozwoli to uniknąć koncentrowania się na dążeniu do sukcesu projektu widzianego tylko z jednego punktu widzenia, np. do unikania jedynie ryzyka niespełnienia wymagań unijnych bądź tylko do unikania ryzyka „kłopotów”, dodatkowych zadań dla pracowników urzędu.

Perspektywy zarządzania ryzykiem, jakie proponujemy dla projektów europejskich to:

- 1) perspektywa społeczności lokalnej,
- 2) perspektywa Unii Europejskiej,
- 3) perspektywa urzędu realizującego projekt.

Cele z punktu widzenia wszystkich trzech perspektyw są różne, a często mogą być nawet sprzeczne. Czym innym jest czas realizacji projektu, jakość projektu czy aspekt funkcjonowania miasta lub gminy w trakcie projektu dla społeczności lokalnej, komisji Unii Europejskiej czy urzędu, który dany projekt realizuje.



Rys. 2. Propozycja perspektyw zarządzania ryzykiem dla projektów unijnych

Perspektywa społeczności lokalnej obejmuje mieszkańców okolicy, w której realizowany jest projekt, użytkowników dróg i mostów, ekologów, przeciwników realizacji projektu, przedsiębiorstwa znajdujące się w okolicy.

W ramach perspektywy urzędu najistotniejszą rolę odgrywają perspektywa czasu, kosztów, wykopalisk, przetargów oraz wykonawców.

W ramach perspektywy Unii Europejskiej najistotniejsze są czas, koszty, parametry techniczne projektu, spełnienie wymagań zawartych w strategii rozwoju, spełnienie celów programów europejskich.

Celem zarządzania ryzykiem projektowym zgodnego z prezentowanym podejściem będzie zapewnienie „sukcesu” projektu widzianego ze wszystkich trzech perspektyw w założonej wielkości (np. 80% – zakłada to mierzalność zapewnienia sukcesu, który problem będzie przedmiotem dalszych badań). Ryzyko nieosiągnięcia sukcesu będzie w trakcie realizacji monitorowane we wszystkich trzech perspektywach. Jeśli zostanie stwierdzone niebezpieczeństwo nieosiągnięcia sukcesu w założonym stopniu, konieczna będzie interwencja i poprawa wskaźników w danej perspektywie. W ramach każdej perspektywy opracowane zostaną cele, które trzeba zrealizować, aby osiągnąć zamierzony efekt, a także zestaw mierników mierzących stopień realizacji celów.

Trzeba podkreślić, iż prezentowane podejście na razie znajduje się w fazie przygotowania. Omawiany materiał to wstępna propozycja podejścia. Kompletna metoda zarządzania ryzykiem zawierać będzie model matematyczny oceny ryzyka projektowego. Poszczególnym perspektywom i związanym z nimi źródłom ryzyka zostaną przyporządkowane wagi obrazujące wpływ na ryzyko projektowe w poszczególnych fazach (planowania, wnioskowania o dofinansowanie, realizacji, kontroli oraz zamknięcia i rozliczenia projektu) oraz na całkowite ryzyko projektów europejskich. W zależności od otrzymanego poziomu ryzyka w poszczególnych perspektywach oraz fazach przygotowane zostaną sposoby reakcji na ryzyko. Najtrudniejszym problemem będzie przedstawienie wielkości jakościowych (np. doświadczenie kierownika projektu, stopień spełnienia warunków konkursu o dofinansowanie, czy nastroje społeczności lokalnej) w postaci mierzalnej tak, aby można było formułować mierzalne cele zarządzania ryzykiem i być w stanie ocenić stopień ich spełnienia w poszczególnych perspektywach.

Dalej przedstawiony zostanie przykład ukończonego już projektu europejskiego, który ma pokazać, że zrównoważone podejście do zarządzania ryzykiem projektów europejskich może być istotne i przydatne.

6. Studium przypadku

Projekt „Przebudowa Mostów Warszawskich we Wrocławiu” to projekt współfinansowany przez Unię Europejską ze środków Europejskiego Funduszu Rozwoju Regionalnego w ramach Zintegrowanego Programu Operacyjnego Rozwoju Regionalnego. projektu „Przebudowa Mostów Warszawskich we Wrocławiu” była rozbudowa i modernizacja infrastruktury drogowo-mostowej poprzez dostosowanie jakości i funkcjonalności ulic do standardów Unii Europejskiej. Zakres projektu obejmował remont starych i budowę nowych Mostów Warszawskich, modernizację jezdni na odcinkach objętych projektem wraz z modernizacją torowiska tramwajowego oraz przebudowę infrastruktury technicznej, a także budowę ekranów akustycznych przy zabudowie mieszkaniowej.

Łączna wartość projektu, początkowo zakładana na poziomie 94 mln zł, ostatecznie osiągnęła wartość ponad 106 mln zł (106 145 449,73 zł). Dofinansowanie projektu ze środków pochodzących z funduszy unijnych zamknęło się na poziomie ponad 56 mln zł (56 034 396,00 zł). Początkowo wartość dofinansowania wynosiła ponad 58 mln zł, co stanowiło 75% wydatków kwalifikowanych w ramach projektu. Realizacja prac projektowych trwała od II kwartału 2006 roku (ogłoszenie o zamówieniu w Urzędzie Zamówień Publicznych zostało opublikowane 29 sierpnia 2005, planowany termin rozpoczęcia prac – I kwartał 2006 do IV kwartału 2008 (pierwotna data zakończenia projektu – II kwartał 2008).

Projekt „Przebudowa Mostów Warszawskich we Wrocławiu” zakończył się z trzymiesięcznym opóźnieniem. Ostateczny termin rozliczenia projektu został wyznaczony na koniec sierpnia 2008. Przekroczenie tego terminu wiązałoby się z utratą unijnej dotacji w wysokości ponad 56 mln zł. Wszelkie koszty związane z realizacją projektu musiałyby wówczas zostać poniesione w całości przez Wrocław.

Podstawowe utrudnienia i problemy, które miały bezpośredni wpływ na opóźnienie projektu i ryzyko utraty dotacji to:

- opóźnienie związane z procedurą przetargową,
- opóźnienie w rozpoczęciu prac projektowych (powiązanie z innymi remontami w centrum Wrocławia),
- opóźnienie związane z nieprzewidzianymi odkryciami w czasie prac ziemnych (wykopaliska).

Opóźnienia związane z procedurą przetargową wynikały przede wszystkim z faktu, iż pierwotna decyzja inwestora o wyborze wykonawcy została oprotestowana, co z kolei wymagało ponownej oceny ofert, a w ostateczności odwołanie się do zespołu arbitrażowego. Dodatkowo rozpoczęcie inwestycji zostało opóźnione przez inwestora ze względu na zazębianie się innych remontów w obrębie centrum Wrocławia. Inwestor uznał, iż uruchomienie projektu w zakładanym terminie spowoduje zbyt duże utrudnienia komunikacyjne, a w efekcie paraliż komunikacyjny. W wyniku opóźnień przetargowych oraz konieczności wstrzymania inwestycji do czasu udrożnienia się alternatywnych ciągów komunikacyjnych (również w przebudowie – Most Szczytnicki i pl. Grunwaldzki) prace remontowe rozpoczęły się z trzymiesięcznym opóźnieniem. Dodatkowe opóźnienia związane z realizacją tego projektu związane były z odkryciem nieprzewidzianych utrudnień w pracach ziemnych i wykopaliskowych. W dniu Odry, w miejscu, gdzie miały być wstawione nowe filary mostu, odkryto jeszcze pozostałości po elementach konstrukcyjnych starego VIII-wiecznego mostu. Elementy nie były zaznaczone w żadnej dokumentacji, a opóźnienie wynikło z faktu, iż stawianie nowych pali mostowych można było rozpocząć dopiero po wydobyciu starych. W związku z przekroczeniem harmonogramu prac oraz ryzykiem utraty całości dotacji wykonawca musiał zatrudnić dodatkowych pracowników oraz zwiększyć wymiar czasu pracy, aby zdążyć w wyznaczonym przez Ministerstwo Rozwoju Regionalnego ostatecznym terminie zakończenia prac oraz rozliczenia projektu.

Zastosowanie podejścia zrównoważonego do zarządzania ryzykiem we wszystkich fazach omówionego projektu złagodziłoby skutki niekorzystnych zdarzeń, które miały miejsce i doprowadziły do opóźnienia projektu, a także problemów związanych z dotrzymaniem nawet tego przesuniętego terminu.

Uwzględnienie od początku w perspektywie urzędu problemu przetargów i wykonawców, a także problemu stresu/przeciążenia pracowników i dodatkowych kosztów pracy, wynikających z konieczności zamknięcia projektu w pewnym określonym terminie, który był na pewno znany z odpowiednim wyprzedzeniem i wynikał z perspektywy Unii Europejskiej, pozwoliłoby zapewne zmniejszyć opóźnienie, a także dotrzymać terminu bez ponoszenia dodatkowych kosztów, finansowych i ludzkich. Podobnie, należało od początku, już na etapie planowania i wnioskowania projektu, brać pod uwagę możliwość nieoczekiwanych znalezisk archeologicznych, gdyż takie ryzyko jest dość powszechne i tak naprawdę przy tego typu projektach nie powinno być niespodzianką. Od początku należało też uwzględniać perspektywę lokalnej społeczności, w tym przypadku mieszkańców miasta i funkcjonowania miasta jako całości. Inne remonty przeprowadzane w mieście nie są żadną tajemnicą i powinny one być automatycznie brane pod uwagę przy zarządzaniu ryzykiem każdego projektu, którego realizacja będzie miała wpływ na funkcjonowanie miasta. Należy jednak wszystkie te perspektywy potraktować w sposób zrównoważony. Mieszkańcy miasta muszą i mogą zaakceptować pewien poziom kłopotów komunikacyjnych, jeśli jest to konieczne z punktu widzenia perspektywy Unii Europejskiej (ryzyko utraty dotacji) czy urzędu (brak wystarczającej liczby zasobów do tego, by z danym remontem poczekać i potem musieć pracować w zwiększonym wymiarze godzin, by projekt dokończyć w wymaganym terminie). Pracownicy urzędu także muszą być gotowi na pewną zwiększoną dyspozycyjność w określonych okresach. Należy także wykorzystać możliwości negocjacji z Unią Europejską, by minimalizując ryzyko niespełnienia jej wymagań, uwzględniać jednocześnie ryzyko w pozostałych perspektywach.

Zakończenie

W opracowaniu przedstawiono koncepcję zrównoważonego podejścia do zarządzania ryzykiem projektów współfinansowanych przez Unię Europejską. Podejście to ma na celu zminimalizowanie ryzyka niepowodzenia projektu obserwowanego z różnych punktów widzenia, jako że sukces tego rodzaju projektów może oznaczać co innego dla mieszkańców miasta, co innego dla Unii Europejskiej, a co innego dla urzędników czy wykonawców zajmujących się jego realizacją. Proponowane podejście wymaga dalszej konkretyzacji i testowania na konkretnych projektach tak, by stało się rzeczywiście przydatnym i akceptowanym narzędziem przyczyniającym się do lepszego wykorzystania dostępnych dla naszego kraju funduszy europejskich.

Literatura

1. Jaruga A. (2000). Zrównoważona Karta Dokonań w systemie zarządzania strategicznego. W: Controlling i rachunkowość zarządcza w firmie. Warszawa, 1.
2. Kaplan R.S., Norton D.P. (2001). Strategiczna Karta Wyników. Jak przełożyć strategię na działanie. Wydawnictwo Naukowe PWN, Warszawa.
3. Litwa P. (2003). Zrównoważona Karta Dokonań – implementacja strategii do rzeczywistości. Zeszyty Naukowe. AE, Kraków, 618.
4. Portal Funduszy Europejskich.
<http://www.funduszeuropejskie.gov.pl/slownik/Strony/Stronaglowna.aspx>
(21.02.2009).
5. Pritchard C.S. (2002). Zarządzanie ryzykiem w projekcie. Teoria i praktyka. WIG-Press, Warszawa.
6. Zarządzanie projektem europejskim. (2007). Red. M. Trocki, B. Gucza. PWE, Warszawa.

Dorota Kuchta

Politechnika Wrocławska

NOWA KONCEPCJA ODPORNOŚCI PLANÓW PROJEKTÓW

Wstęp

Każdy projekt jest z definicji przedsięwzięciem w pewnym sensie unikatowym. Z tego względu, a także z powodu zazwyczaj zmiennego otoczenia, w jakim jest realizowany, jest narażony na ryzyko (rozumiane jako zdarzenia, o których wiemy, że mogą mieć miejsce, i że jeśli tak będzie, mogą mieć znaczący negatywny wpływ na pomyślną realizację projektu), a jego planowanie, jak również realizacja, odbywa się w warunkach niepewności i niepełnej wiedzy.

Wszystkie te czynniki sprawiają, że zarządzający projektami borykają się z różnego typu problemami. Stanem idealnym, czy przynajmniej najbardziej pożądanym, byłaby zapewne sytuacja, kiedy projekty przebiegałyby ściśle według planu. Wtedy firma jako całość mogłaby lepiej planować całą swoją działalność, być bardziej wiarygodna wobec klientów, aktualnych i przyszłych, a i pracownicy nie musieliby przeżywać „gorących okresów”, kiedy trzeba pracować dzień i noc, bo stało się coś, czego w planie nie przewidziano, a ustalone z klientem terminy zakończenia projektów gonią. Niestety, chyba żaden realny projekt nie został zrealizowany dokładnie według planu, a doniesienia z praktyki mówią o bardzo znacznych odstępstwach, wywołujących różnego rodzaju skutki w różnych aspektach danego projektu i jego otoczenia. Trochę mniej ambitnym celem byłoby wymaganie, by każdy projekt realizowany przez firmę kończył się w taki sposób, by można to było uznać za sukces, czyli by najważniejsze parametry końcowego produktu projektu były zgodne z oczekiwaniami. W takim przypadku realizacja planu jako takiego by nas tak bardzo nie interesowała, tylko sam produkt końcowy. Jest to podejście dość dyskusyjne, chociażby dlatego, że na pewno nie ułatwia planowania pracy firmie jako takiej, ale są zwolennicy tezy, że tak naprawdę interesuje nas tylko efekt końcowy projektu, a do planu trzeba podchodzić w sposób elastyczny [4]. Niestety, praktyka pokazuje, że i ten mniej ambitny cel rzadko jest realizowany.

Można by powiedzieć, że tak już jest, że projekty są czymś zmiennym i zawsze takimi będą, a w dzisiejszym pełnym zmienności świecie może nawet będą nabierały tej cechy w coraz większym stopniu. Nie da się zupełnie wyeliminować niepewności i zmienności związanych z planowaniem i realizacją projektów. Niemniej jednak jest często możliwość stworzenia takiego planu projektu, który będzie odporny na ryzyka, czyli takiego, że nawet jeśli będą miały miejsce różne zdarzenia mogące doprowadzić do niekorzystnych konsekwencji dla projektu, to projekt się nie „zawali” – w takim sensie, w jakim jest to dla nas istotne: czyli albo nie zawali się plan i będzie dość wiernie realizowany, albo „przynajmniej” efekt końcowy będzie zadowalający. Odporność może w zasadzie dotyczyć tylko ryzyka (w sensie podanej wyżej definicji), czyli zdarzeń, które na etapie planowania możemy przewidzieć, przed zdarzeniami zupełnie nieprzewidywalnymi chronić się trudniej, niemniej jednak już stworzenie planu odpornego na zidentyfikowane zdarzenia (przy założeniu, że identyfikacja tych zdarzeń [3; 9] przeprowadzona została rzetelnie) to bardzo duży krok w kierunku stanu pożądanego.

W niniejszym opracowaniu przedstawimy znane z literatury podejścia do odporności planu projektu, jako że pojęcie to można rozumieć w różny sposób, a następnie propozycję nowej koncepcji, która zostanie zilustrowana przykładem.

1. Różne koncepcje odporności planu projektu

Odporny plan projektu jest zazwyczaj definiowany jako taki, który będzie się niewiele różnić od rzeczywistej realizacji projektu. Kwestią otwartą jest znaczenie wyrażenia „niewiele się różnić” w odniesieniu do planu projektu. W literaturze [5; 11] wyróżnia się odporność jakościową, która po prostu oznacza, że rzeczywista wartość wybranej charakterystyki projektu (najczęściej czasu trwania) będzie nie gorsza lub niewiele gorsza od planowanej (to podejście odpowiada opisanemu we wstępie „mniej ambitnemu” celowi, związanemu „jedynie” z efektem końcowym projektu) oraz harmonogramową, która oznacza, że rzeczywiste wartości z pewnego zbioru charakterystyk planu (np. zbiór momentów rozpoczęcia zadań) nie będą się wiele różniły od planowanych (to podejście odpowiada przedstawionemu we wstępie dążeniu do możliwie wiernej realizacji planów projektów).

Narzędzia pozwalające zapewnić tak rozumianą odporność to między innymi:

1. Rozważanie różnych scenariuszy i planowanie projektu przy pesymistycznych czy bliskich pesymistycznym założeniach [1; 7; 8].

2. Rozważanie rozkładów prawdopodobieństwa poszczególnych parametrów projektu i planowanie przy minimalizacji wartości oczekiwanej różnic między planem a wartościami rzeczywistymi [5].

3. Wstawianie w planie buforów mających chronić przed niespodziewanymi zmianami to, na czego ochronie najbardziej nam zależy (np. czasu realizacji projektu, jeśli interesuje nas odporność jakościowa, czasu rozpoczęcia poszczególnych zadań, jeśli interesuje nas odporność harmonogramowa itp.). Problem rozmieszczenia, a zwłaszcza wielkości buforów nie ma jednoznacznego rozwiązania i zależy od wielu czynników, np. od tego, czy ryzyko związane z poszczególnymi zadaniami zależy od ich długości, czy od innych rzeczy. Stosowanie buforów jest między innymi przedstawione w [4] (odporność jakościowa), [6] (odporność harmonogramowa), [2] (przegląd różnych przypadków stosowania buforów).

W kolejnym rozdziale zostanie przedstawiona inna koncepcja odporności planu projektu, oparta na ogólnej koncepcji systemów odpornych na awarię. W tej koncepcji nie będzie wymagana ani analiza scenariuszy, ani rozważanie rozkładów prawdopodobieństwa, ani stosowanie buforów, choć wszystkie te techniki będą mogły być połączone z proponowaną poniżej, by dać w rezultacie tak odporny plan projektu, jakiego decydent będzie potrzebował. Celem nowej koncepcji jest przede wszystkim wspomaganie zarządzających projektami w osiągnięciu odporności jakościowej.

2. Nowa koncepcja odporności planu projektu

W [10] zaproponowana została całościowa koncepcja systemów odpornych na awarie, obejmująca takie systemy, jak ratownicze, w których awaria oznacza nieudzielenie pomocy w sytuacji takiej pomocy bezwzględnie wymagającej. Koncepcja ta obejmuje kilka aspektów, które mogą być przeniesione na zarządzanie projektami. Samo pojęcie awarii, która w systemie ratowniczym oznacza nieudzielenie pomocy, w projekcie będzie miało swój odpowiednik w zdarzeniu, którego skutkiem będzie znaczące odchylenie rzeczywistego wyniku realizacji projektu od wyniku pożądanego. Odporny system ratowniczy to taki, który nie zawiedzie w żadnej rzeczywiście poważnej sytuacji i doprowadzi do udzielenia niezbędnej pomocy, nawet jeśli nie zaraz po zgłoszeniu, to w takim czasie, że pomoc ta jeszcze będzie możliwa. Plan projektu odporny na awarie to będzie – przez analogię – taki plan, który podczas swojej realizacji umożliwi reakcję na wszystkie awarie, czyli zdarzenia mogące mieć znaczący negatywny wpływ na efekt projektu – reakcję, nawet jeśli nie natychmiastową to taką, która w porę uchroni wynik projekt przed znaczącymi odchyleniami od stanu pożądanego.

Poszczególne elementy koncepcji odpornych systemów zaproponowanej w [10] to:

1. Pełna integracja systemu, rozumiana jako świadomość i uwzględnianie wzajemnych zależności między elementami systemu. W przypadku projektów oznacza to identyfikację wszelkich zależności między elementami projektu: między poszczególnymi zadaniami, zasobami, udziałowcami w szerokim tego słowa znaczeniu itp. i nie będzie tu chodziło tylko o zależności typu relacje poprzedzania, które w zarządzaniu projektami są uwzględniane „od zawsze”. Chodzi tu o zależności często ukryte, nieformalne, typu „ktoś ma na coś decydujący wpływ, jest w stanie coś skutecznie zablokować lub przyspieszyć”, „takie a takie zadanie będzie wykonywane lepiej, jeśli inne zadanie będzie wykonane wcześniej, bo będzie miało miejsce zjawisko uczenia się” itp.

2. Stosowanie tzw. procesu decyzyjnego Boyda, który polega na tym, że cykl „zaobserwowanie problemu, zorientowanie się w jego naturze, decyzja i akcja” mają następować jak najszybciej. W przeciwnym przypadku problemy się kumulują i system się blokuje. W przypadku projektów oznaczać to będzie zasadę, że na każdy sygnał należy reagować zaraz, a nie stosować zasadę „zobaczy się później”, nawet jeśli interesuje nas tylko odporność jakościowa, czyli ostateczny rezultat projektu.

3. Wspólny sposób myślenia zespołu, rozumiany jako umiejętność przewidywania reakcji kolegów z zespołu. Oczywiście jest, że dotyczy to w szczególności członków ekip ratowniczych, ale i członków zespołu projektowego powinno łączyć podobne podejście do projektu i sposobu jego realizacji. Tutaj szczególnie ważna jest kultura przedsiębiorstwa, której wagę podkreślił w kontekście zarządzania projektami już Godratt [4].

4. Standardowe procedury postępowania. Wiele przedsiębiorstw wypracowało takie procedury na niemal każdą przewidywalną sytuację (w [10] jako koronny przykład wspomniana jest Toyota). W przypadku zarządzania projektami takie procedury są również bardzo istotne, zwłaszcza w kontekście zarządzania ryzykiem, i wiele firm je wypracowało [3]. Wymuszają one na zespole projektowym systematyczne podejście do zarządzania ryzykiem, nie pozwalają ulec zwodniczej magii zasady „jakoś to będzie”, która zazwyczaj prowadzi do poważnych problemów związanych z realizacją celów projektu.

5. Odporne zaprojektowanie systemu, czyli takie zaplanowanie jego działania, by reagował on na nadchodzące sygnały w sposób jak najbardziej skuteczny. W przypadku projektów będzie tu chodziło o odpowiednie zaplanowanie projektu.

Właśnie próbą przeniesienia tego aspektu odpornych systemów na zarządzanie projektami zajmiemy się w dalszym ciągu nieco bliżej.

Według [10], system zaprojektowany w sposób odporny to system, który spełnia między innymi następujące warunki:

1. Różnicuje czy raczej hierarchizuje przychodzące sygnały, biorąc od uwagę dwa kryteria:

- krytyczność sygnału, innym słowy wagę sygnału z punktu widzenia jego skutków (inną wagę będzie miał sygnał informujący o zdarzeniu, którego konsekwencją może być śmierć czy trwałe kalectwo człowieka, inną wagę sygnał informujący o zdarzeniu, którego konsekwencją może być niedyspozycja usuwalna (np. złamanie nogi) czy skutki nie dotyczące życia i zdrowia ludzkiego,
- „czas do katastrofy”, czyli czas mierzony od momentu otrzymania sygnału, za jaki najprawdopodobniej wystąpią spodziewane skutki zdarzenia (np. śmierć).

Ostateczna hierarchia sygnałów i sposób reakcji na nie jest wypadkową tych dwóch kryteriów.

2. Zapewnia, że osoby pracujące w systemie w żadnym momencie nie będą zmuszone do reakcji na silniejszy bodziec, niż są do tego zdolne, przy czym pod pojęciem bodźca rozumiemy sumę sygnałów, na które trzeba w danym momencie zareagować, przemnożonych przez ich wagi. Spełnienie tego warunku oznacza umiejętność oszacowania wielkości bodźców w każdym momencie i zapewnienia dostatecznej liczby zasobów, zapewniającej odpowiednią reakcję zarówno na przeciętne, jak i nadzwyczajne wielkości bodźców.

Powyższe warunki, przeniesione na planowanie projektów ukierunkowane na skuteczne zarządzanie ryzykiem, oznaczają:

3. Uwzględnianie zarówno skutków, jakie poszczególne zdarzenia mogą wyrzucić na ostateczny efekt projektu, jak i czasu, jaki jeszcze pozostał do momentu, kiedy skutki te będą nieodwracalne. Oznaczać to może np. planowanie zadań, przy realizacji których mogą wystąpić zdarzenia mające bardzo negatywne skutki dla rezultatu projektu, wcześniej, by jeszcze był czas na ewentualne naprawienie szkód.

4. Takie rozplanowanie zadań, by w żadnym momencie obciążenie zespołu projektowego związane z zapobieganiem i reakcją na niekorzystne zdarzenia nie przerosło jego możliwości.

Teraz proponujemy model matematyczny, który jest próbą formalnego uwzględnienia proponowanego podejścia w zarządzaniu projektami.

Rozpatrujemy projekt złożony z n zadań. Dla każdego zadania jest dany planowany czas trwania $d_i (i=1, \dots, n)$. Ponadto znany jest zbiór $S = \{(i_1, i_2) : i_1, i_2 \in \{1, \dots, n\}, i_1 < i_2\}$ taki, że zadanie i_1 poprzedza zadanie i_2 (przy czym poprzedzanie rozumiemy w ten sposób, że koniec zadania i_1 musi nastąpić przed początkiem zadania i_2). Znany jest też horyzont H , czyli moment, kiedy projekt bezwzględnie musi być zakończony (zakładamy, że roz-

poczynamy go w na początku 1. jednostki czasu, np. miesiąca czy tygodnia, wówczas H jest numerem ostatniej jednostki czasu okresu, w którym w ogóle można rozpatrywać realizację projektu). Naszym celem jest minimalizacja momentu zakończenia całego projektu, przy uwzględnieniu następujących ograniczeń:

- relacji poprzedzania,
- wymogu zapewnienia zdolności zespołu projektowego do sprostania w każdym momencie realizacji projektu obciążeniom wynikającym z zarządzania ryzykiem (czyli konieczności kontroli ryzyka, reakcji na niepokojące sygnały, zapobiegania niekorzystnym zdarzeniom, zmniejszania skutków niekorzystnych zdarzeń itp.).

Jeśli chodzi o drugi warunek zakładamy, że dana jest funkcja $L(t), t = 1, \dots, H$, obrazująca zdolności zespołu projektowego w poszczególnych momentach do realizacji zadań związanych z ryzykiem. Funkcja ta może być stała, ale wydaje się, że bardziej realistycznym założeniem jest przyjęcie funkcji rosnącej: doświadczenie pokazuje, że im bliżej ostatecznego terminu zakończenia projektu, tym zespół projektowy jest bardziej zmobilizowany i zdolny do większego zaangażowania, podczas gdy na początku realizacji projektu znaczną rolę odgrywa zazwyczaj „syndrom studenta” [4], zgodnie z którym zespół pewną część zadań do wykonania odkłada „na jutro”. Problemem pozostaje kwantyfikacja zdolności $L(t), t = 1, \dots, H$, co będzie przedmiotem dalszych badań.

Ten sam sposób kwantyfikacji będzie przyjęty przy określaniu siły bodźca płynącej z poszczególnych zadań, czyli wielkości problemów, jakie to zadanie może potencjalnie sprawić zespołowi projektowemu z tytułu zarządzania ryzykiem. Siła bodźca nie będzie określona stałą, lecz również funkcją $t = 1, \dots, H$, i to zawsze funkcją niemalejącą, zazwyczaj dodatkowo kawałkami ściśle rosnącą. Wynika to z tego, iż próbujemy uwzględnić nie tylko wielkość potencjalnych skutków zdarzeń związanych z danym zadaniem z punktu widzenia ostatecznego rezultatu projektu, ale także „czas do katastrofy”, czyli czas, jaki dzieli moment realizacji danego zadania od momentu zakończenia projektu. Im zadanie będzie zaplanowane bliżej końca projektu, tym bodźce związane z tym zadaniem, obciążenie wynikające z konieczności zarządzania ryzykiem związanym z tym zadaniem będzie większe, bo będzie mniej czasu na ewentualne naprawienie szkód i trzeba będzie reagować zdecydowanie i intensywnie. Dlatego przyjmujemy, że dla i -tego ($i = 1, \dots, n$) zadania dana jest niemalejąca funkcja $R_i(t), t = 1, \dots, H$, reprezentująca potencjalny wysiłek, jaki będzie wymagany od zespołu projektowego ze względu na zarządzanie ryzykiem związanym z i -tym zadaniem, jeśli zadanie to będzie realizowane w t -tej jednostce czasu.

Zmiennymi decyzyjnymi modelu będą zmienne binarne $x_i^t, i=1, \dots, n; t=1, \dots, H-1$, które będą przyjmować wartość 1 wtedy i tylko wtedy kiedy i -te zadanie rozpocznie się w t -tej jednostce czasu. Minimalizowaną funkcją celu, reprezentującą moment zakończenia projektu, będzie funkcja $F(x_i^t, i=1, \dots, n; t=1, \dots, H-1) = \max_{i=1, \dots, n} \sum_{t=1}^{H-1} t \cdot (x_i^t + d_i)$. Ograniczenia modelu będą następujące

$$\sum_{t=1}^{H-1} t(x_{i_1}^t + d_{i_1}) \leq \sum_{t=1}^{H-1} t x_{i_2}^t \text{ dla każdego } (i_1, i_2) \in S \quad (1)$$

$$\sum_{t=1}^{H-1} x_i^t = 1 \text{ dla każdego } i=1, \dots, n \quad (2)$$

$$F(x_i^t, i=1, \dots, n; t=1, \dots, H-1) \leq H \quad (3)$$

$$\sum_{i=1}^n x_i^t R_i(t) + \sum_{i=1}^n \sum_{\substack{s=t-d_i, \\ s \geq 1}}^{t-1} x_i^s R_i(t) \leq L(t) \text{ dla każdego } t=1, \dots, H \quad (4)$$

Ograniczenie (1) reprezentuje relacje poprzedzania, ograniczenie (2) oczywisty fakt, że każde zadanie może i musi się rozpocząć tylko raz, ograniczenie (3) zapewnia, że z realizacją projektu zmieścimy się w zadanym horyzoncie. Natomiast ograniczenie (4) po lewej stronie sumuje wysiłek związany z zarządzaniem ryzykiem, a wynikający z zadań realizowanych w poszczególnych jednostkach czasu (pierwszy składnik sumuje odpowiednie wartości dla zadań rozpoczętych w t -tej jednostce czasu, a drugi składnik dla zadań rozpoczętych wcześniej, ale jeszcze trwających w t -tej jednostce czasu) i zapewnia, że nie będzie on większy niż możliwości zespołu projektowego w danej jednostce czasu.

Zaproponowaną koncepcję można w prosty sposób rozszerzyć na przypadek projektów z bardziej zróżnicowanymi formami relacji poprzedzania, z narzuconymi limitami na zasoby itp. Tak naprawdę można powiedzieć, że w proponowanej koncepcji funkcja $L(t), t=1, \dots, H$ może być interpretowana jako dostępność pewnego zasobu, a funkcje $R_i(t), t=1, \dots, H, i=1, \dots, n$ jako zapotrzebowania zadań na ten zasób. Jednak zapotrzebowanie poszczególnych zadań na ten zasób jest innego rodzaju niż w przypadku zasobów w sensie stricte: zależy on również od momentu, w którym zadanie będzie wykonywane, i może się zwiększać wraz z przesuwaniem zadania w planie w kierunku końca projektu.

W następnym rozdziale przedstawimy prosty przykład, ilustrujący zaproponowaną koncepcję.

3. Przykład liczbowy

Rozważamy projekt złożony z pięciu zadań ($n=5$). Mamy:

- $d_1 = 1, d_2 = 2, d_3 = 1, d_4 = 2, d_5 = 3$
- $S = \{(2,4), (1,5), (2,5), (2,3)\}, H = 10$

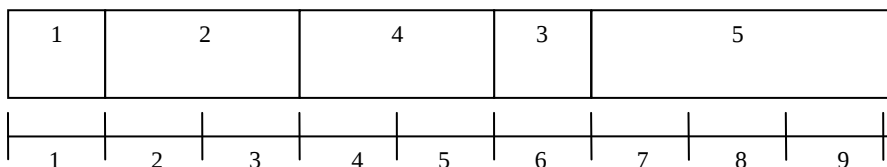
Rozważmy następujące przypadki:

Tabela 1

Dane do przykładu liczbowego, przypadek a)

t	1	2	3	4	5	6	7	8	9	10
$L(t)$	5	6	7	8	9	9	10	10	10	10
$R_1(t)$	5	6	7	8	9	9	10	10	10	10
$R_2(t)$	5	6	7	8	9	9	10	10	10	10
$R_3(t)$	3	4	5	6	7	8	9	10	10	10
$R_4(t)$	4	4	6	8	9	10	10	10	10	10
$R_5(t)$	5	6	7	8	9	9	10	10	10	10

Przy takich danych rozwiązaniem sformułowanego w poprzednim rozdziale zadania (można je wyznaczyć np. za pomocą SOLVERA w programie EXCEL) będzie następujący harmonogram wykonywania zadań:



Rys. 1. Optymalny plan projektu w sytuacji a)

Widać zatem, że projekt nie może być ukończony szybciej niż po 9 jednostkach czasu (mimo iż ścieżka krytyczna projektu nie uwzględniająca ryzyka i potrzeb oraz ograniczeń wynikających z zarządzania nim ma długość 7). Przy wartościach z tabeli 1 żadne dwa zadania projektu nie mogą być wykonywane równolegle, bo przy równoległym wykonywaniu któregośkolwiek z zadań mielibyśmy do czynienia z planem nieodpornym (zespół projektowy potencjalnie nie

byłby w stanie sobie poradzić z potrzebami wynikającymi z zarządzania ryzykiem). Przy tego typu problemach zazwyczaj istnieją alternatywne plany optymalne, w naszym przypadku takim planem byłby plan, przy którym kolejność wykonywania zadań 1 i 2 byłaby odwrócona w stosunku do rys. 1. Ale już kolejności wykonywania zadań 3 i 4 nie można by odwrócić: wprowadzie relacje poprzedzania na to by zezwalały, ale zadanie 4 musi być najpóźniej ukończone w 5 jednostce czasu ze względu na dane z tabeli 1 – później, bliżej końca projektu, ma ono na tyle wysoką wagę z punktu widzenia ryzyka, że zespół projektowy już mógłby nie zdążyć poradzić sobie z konsekwencjami niekorzystnych wydarzeń, do jakich to zadanie może doprowadzić.

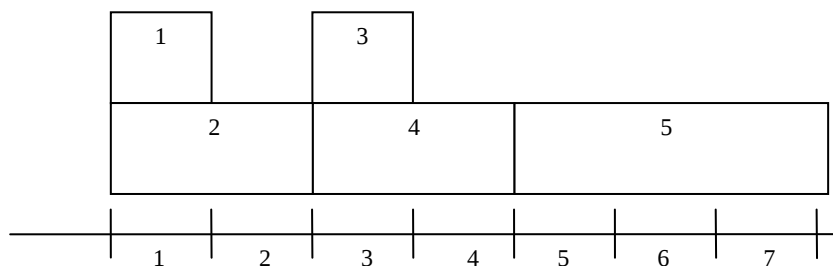
Rozpatrzmy teraz inną sytuację (dla tego samego przykładu):

Tabela 2

Dane do przykładu liczbowego, przypadek b)

T	1	2	3	4	5	6	7	8	9	10
$L(t)$	9	9	9	10	10	10	10	10	10	10
$R_1(t)$	2	3	4	5	6	7	9	9	10	10
$R_2(t)$	5	6	7	8	9	9	10	10	10	10
$R_3(t)$	2	3	4	5	6	7	9	9	10	10
$R_4(t)$	5	5	5	8	9	10	10	10	10	10
$R_5(t)$	5	6	7	8	9	9	10	10	10	10

W tym przypadku okazuje się, że projekt uda się zrobić w czasie równym długości ścieżki krytycznej w przypadku nie uwzględniającym zarządzania ryzykiem:



Rys. 2. Optymalny plan projektu w sytuacji b)

Warto, jeśli jest to możliwe, zająć się zadaniami, które im później będą realizowane, tym bardziej mogą w nieodwracalny sposób zaszkodzić projektowi, możliwie wcześniej (sytuacja b), zadania 1 i 3), ale nie można planować takiego postępowania, jeśli jest ono nierealne i oznacza nadmierne przeciążenie zespołu projektowego (sytuacja a). Inaczej cały system mógłby się zablokować i projekt i tak by się przedłużył, tyle że w sposób niezaplanowany.

Podsumowanie

W pracy wykorzystano koncepcję odpornych systemów (głównie ratowniczych) jako podstawę koncepcji odpornego planowania projektów, gdzie odporność jest rozumiana w tym sensie, że niezależnie (wyłączając całkowicie nieprzewidywalne zdarzenia) od tego, co się zdarzy podczas realizacji projektu, najważniejsze jego cele zostaną osiągnięte. Dalszych badań wymaga problem ilościowego wyrażenia potrzeb związanych z zarządzaniem ryzykiem z jednej strony i możliwości zespołu projektowego w tym zakresie z drugiej. Warto będzie również rozważyć hybrydowe podejścia do odporności planu projektu: wykorzystujące zarówno podejście zaproponowane w niniejszym opracowaniu, jak i inne podejścia, już znane z literatury.

Literatura

1. Aissi H., Bazgan C., Vanderpooten V. (2005). Complexity of the Min-max and Min-max Regret Assignment Problems. *Operations Research Letters*, 33, 634-640.
2. Chtourou H. Haouari M. (2008). A Two-stage-priority-rule-based Algorithm for Robust resource-constrained Project Scheduling. *Computers & Industrial Engineering* 55(1), 183-194.
3. Courtot H. (1998). La gestion des risques dans les projets. *Economica*, Paryż.
4. Goldratt E.M. (2000). Łańcuch krytyczny. Werbel.
5. Herroelen W., Leus R. (2004). Project Scheduling under Uncertainty: Survey and Research Problems. *European Journal of Operational Research*, 165(2), 289-306.
6. Kouvelis P., Yu G. (1997). Robust Discrete Optimization and its Applications. Kluwer Academic Publishers, Boston.
7. Kuchta D., Kobylański P. (2007). A Note on the Paper by M. A. Al-Fawzan and M. Haouari About a Bi-objective Problem for Robust Resource-constrained Project Scheduling. *International Journal of Production Economics*, 107, 2, 496-501.

-
8. Lebedev V., Averbakh I. (2006). Complexity of Minimizing the Total Flow Time with Interval Data and Minmax Regret Criterion. *Discrete Applied Mathematics*, 154, 2167-2177.
 9. Pritchard C. (2002). Zarządzanie ryzykiem w projekcie. Teoria i praktyka. WIG-Press, Warszawa.
 10. Véronneau S., Cimon Y. (2007). Maintaining Robust Decision Capabilities: An Integrative Human – Systems Approach. *Decision Support Systems*, 43(1), 127-140.
 11. Vonder V. de, Demeulemeester E., Herroelen W. (2008). Proactive Heuristic Procedures For Robust Project Scheduling: An Experimental Analysis. *European Journal of Operational Research*, 189(3), 723-733.

Dorota Kuchta

Ewa Ptaszyńska

Politechnika Wrocławska

ZARZĄDZANIE RYZYKIEM POPRZEC UCZENIE SIĘ W PROJEKTACH EUROPEJSKICH

Wstęp

Zarządzanie ryzykiem jest ogólnie uważane za jeden z najważniejszych elementów zarządzania każdym projektem, jako że każdy projekt jest z definicji w znacznym stopniu unikatowy, niepowtarzalny, realizowany w niestabilnym środowisku. Firmy, które realizują wiele projektów i mają w swojej branży wieloletnie doświadczenie, jakoś sobie z zarządzaniem ryzykiem radzą, choć nie do końca, bo odsetek projektów biznesowych, które kończą się niepowodzeniem, jest jednak ogromny. Ale w znacznie gorszej sytuacji są organizacje, które takiego doświadczenia nie mają. Takim organizacjami są np. urzędy, które stanęły wobec konieczności po pierwsze, starania się o fundusze europejskie, a po drugie, realizacji i rozliczania takich projektów. Organizacje te nie mają wielkiego doświadczenia w tym zakresie, ich styl pracy był też zawsze z natury swojej inny niż styl pracy osób zatrudnionych w firmach walczących o utrzymanie się na rynku i zyski. Do tego wszystkiego dochodzi fakt, że realizacja projektów europejskich „obchodzi” znacznie szersze grono osób niż realizacja projektu biznesowego. Są to zazwyczaj projekty ważne dla miejscowej społeczności, ich realizacja odbywa się na otwartym terenie (budowa mostu, szkoły), jest poddana bardzo szerokiej ocenie, mniej lub bardziej emocjonalnej i mniej lub bardziej fachowej. Do tego sposób rozliczania projektów europejskich jest też specyficzny, powiązany z daleką Brukselą, a każde opóźnienie może skutkować utratą przyznanych funduszy. Dlatego zarządzanie ryzykiem w takim projektach jest po pierwsze, bardzo ważne, po drugie, musi być inne niż w projektach biznesowych, po trzecie, wymaga rozwijania odpowiednich metod, bo jest jeszcze – można powiedzieć - w powijkach.

Niniejsze opracowanie proponuje system zarządzania ryzykiem projektów europejskich, który z samego swojego założenia chce traktować problem kompleksowo, czyli niejako wymusić korzystanie z doświadczeń, uczenie się na błędach popełnionych przez własną i inne organizacje, które takim projektami

się zajmują. Każde inne podejście będzie skazane na niepowodzenie lub jego wykorzystanie zajmie niewspółmiernie więcej czasu, powodując, że przy kolejnych projektach europejskich występować będą kolejne poważne problemy, irytujące lokalną społeczność, których można by uniknąć, gdyby istniał system pozwalający na uczenie się i korzystanie z doświadczeń swoich i innych.

1. Zarządzanie ryzykiem projektu a uczenie się

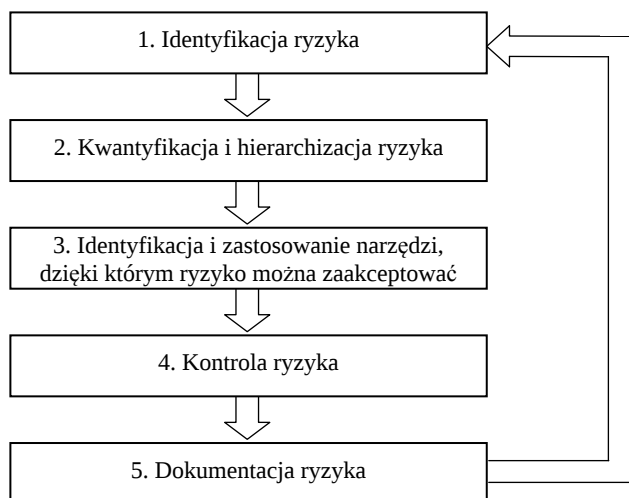
Zagadnieniu ryzyka projektu poświęconych jest wiele pozycji literaturowych, które w ogólnym zarysie spojrzenia na ryzyko projektu i zarządzanie nim są ze sobą zgodne. Literatura zawiera również różne definicje ryzyka. Tutaj oparto się na definicji ryzyka projektu podanej przez Courtot [1], którą przyjęto w nieco zmodyfikowanej formie.

Definicja 1

Ryzyko projektowe to zdarzenie, którego wystąpienie jest możliwe i które swoim wystąpieniem stworzy możliwość, że po zrealizowaniu projektu jego data zakończenia, koszt i produkt końcowy nie będą zgodne z planem, a odchylenia od planu w przynajmniej jednym z tych trzech aspektów będą trudne czy wręcz niemożliwe do zaakceptowania.

Przy takiej definicji ryzyka projektowego uzasadnione jest używanie formy liczby mnogiej „rodzaje ryzyka projektu”.

Następnie Courtot podaje poszczególne kroki zarządzania ryzykiem projektu oraz ich opis. Ogólny schemat [1] przedstawia rys. 1.



Rys. 1. Klasyczny proces zarządzania ryzykiem projektu

Źródło: [1].

Celem identyfikacji ryzyka jest sporządzenie listy ryzyka, jakie może dotyczyć planowanego czy już realizowanego projektu. Oznaczmy poszczególne zidentyfikowane rodzaje ryzyka, jako R_i ($i = 1, \dots, n$), gdzie n jest liczbą zidentyfikowanego ryzyka. Następnie, podczas kwantyfikacji i hierarchizacji ryzyka, należy każde zidentyfikowane powiązać z pewnymi liczbowymi charakterystykami. Courtot przyjmuje konieczność wyznaczenia dla każdego ryzyka R_i ($i = 1, \dots, n$) trzech charakterystyk:

- prawdopodobieństwa wystąpienia p_i ,
- wielkości negatywnych konsekwencji wystąpienia ryzyka k_i (w sensie wagi odchyień wspomnianych w definicji 1, do których może doprowadzić i-te ryzyko),
- stopnia trudności wcześniejszego wykrycia przyszłego wystąpienia ryzyka w_i (czyli jakichś sygnałów ostrzegawczych, tak jak zachmurzone niebo ostrzega przed nadchodzącym deszczem).

Wartości p_i, k_i, w_i ($i = 1, \dots, n$) przyjmują wartości z przyjętej skali, najwyższa wartość jest za każdym razem najmniej korzystna dla projektu: przy najwyższych możliwych wartościach p_i, k_i, w_i ($i = 1, \dots, n$) ryzyko R_i jest wysoce prawdopodobne, jeśli wystąpi, będzie miało bardzo duże negatywne konsekwencje dla projektu i dodatkowo nie będzie możliwe zapobiec mu wcześniej, bo pojawi się bez sygnałów ostrzegawczych.

Na podstawie wartości p_i, k_i, w_i ($i = 1, \dots, n$) ustala się hierarchię ryzyka, odrzucając to, które można pominąć w dalszych etapach, ze względu na nikłe prawdopodobieństwo wystąpienia, niewielkie konsekwencje i łatwą wczesną wykrywalność. Niemniej jednak trzeba pamiętać o tym, że znaczenie wyrażen „nikłe prawdopodobieństwo”, „niewielkie konsekwencje” i „wczesna wykrywalność” jest inne w każdym projekcie, dlatego nie można podać gotowego, ogólnego wzoru wyznaczającego hierarchię ryzyka. Będzie ona wyznaczana w sposób specyficzny dla każdego projektu.

W następnym etapie należy zidentyfikować i zastosować narzędzia, dzięki którym zidentyfikowane i nieodrzucone w poprzednim kroku ryzyko stanie się akceptowalne. Przykładem takiego narzędzia jest ubezpieczenie, zmiana formy umowy czy sposobu rozliczania z kontrahentem, przeprowadzenie szkoleń zespołu projektowego itp. Niektórych z tych narzędzi nie można zastosować wcześniej, ale trzeba ich użycie zaplanować na wypadek wystąpienia danego ryzyka.

Podczas realizacji projektu trzeba ryzyko kontrolować, sprawdzając, jakie ryzyko wystąpiło, czy można/trzeba podjąć w związku z tym jakieś kroki, jakie sygnały ostrzegawcze występują, jakie ryzyko musimy w dalszym ciągu brać pod uwagę, a jakie już jest nieaktualne itp.

Ostatnim krokiem jest dokumentacja ryzyka, tego, jakie rodzaje ryzyka wystąpiły, jakie miały konsekwencje, jakie sygnały ostrzegawcze, dzięki jakim krokom udało się niektórych uniknąć, konsekwencje innych zmniejszyć itp. Dokumentacja ryzyka ma być wykorzystywana w kolejnych projektach.

Choć w zasadzie kroki 1, 2, 3 są wykonywane przede wszystkim w fazie planowania projektu, krok 4 w fazie jego realizacji, a krok 5 po zakończeniu projektu, to tak naprawdę cały cykl powinien być powtarzany systematycznie podczas realizacji projektu, choć kroki 1, 2, 3 i 5 będą wtedy – ze względu głównie na brak czasu – realizowane w postaci okrojonej.

To, że uczenie się odgrywa w procesie zarządzania ryzykiem projektu ważną rolę, widać wprost w kroku 5. Celem tego kroku jest wyłącznie zapewnienie możliwości uczenia się na zrealizowanych projektach po to, by można było lepiej zarządzać kolejnymi projektami. Tak naprawdę, choć na rys. 1 produkt kroku 5 stanowi wejście tylko dla kroku 1, wiedza dotycząca już zrealizowanych projektów i związanych z nimi ryzykiem jest istotna również w krokach 2, 3 i 4. Dla każdego z tych kroków można w literaturze znaleźć pewne metody, mniej lub bardziej formalne [1], ale zawsze podkreślana jest zasadnicza rola, jaką we wszystkich tych krokach odgrywa wiedza ekspercka. Wiedza ta może być z kolei zbierana różnymi metodami, np. za pomocą burzy mózgów czy wywiadów, niemniej jednak wydaje się, że najlepszym sposobem udostępnienia wiedzy opartej na doświadczeniu jest sporządzanie odpowiedniej dokumentacji i korzystanie z niej. Trzeba bowiem wziąć pod uwagę fakt, że pracownicy zmieniają miejsce zatrudnienia i często osoby uczestniczące w projektach zabierają wiedzę ekspercką ze sobą. Nie jest ona wtedy dostępna ani poprzez wywiady, ani burze mózgów, i nowe zespoły projektowe muszą zaczynać od nowa, ucząc się często na własnych błędach. Co więcej, istnienie odpowiedniej pisemnej informacji umożliwiłoby również korzystanie z doświadczeń innych firm, realizujących projekty podobnego typu. Bo wprawdzie każdy projekt jest z definicji unikatowy, jednak występują grupy projektów o podobnych charakterze, gdzie doświadczenie, historia zarządzania ryzykiem miałyby ogromne znaczenie dla coraz lepszego zarządzania kolejnymi projektami.

Problem polega na tym, że sporządzenie odpowiedniej dokumentacji, czytelnej dla osób realizujących inne projekty w przyszłości, jest bardzo pracochłonne i trudne. Rozmowy z praktykami z polskich firm wskazują jasno, że z braku czasu, a także jasnych wskazówek co do formy zapisu informacji, krok 5 realizowany jest w praktyce rzadko. Mamy tu wyraźny rozdzźwięk między teorią i praktyką.

Można temu zapobiec tylko poprzez tworzenie jasnych procedur sporządzania odpowiedniej dokumentacji, dostosowanych do projektów poszczególnych typów. Inaczej będzie dokumentowane ryzyko projektów wysłania statku kosmicznego, inaczej wytworzenia i wdrożenia oprogramowania. Musi też

zostać sformalizowany proces korzystania z tej dokumentacji, czyli – upraszczając – strzałka wychodząca od p. 5 na rys. 1 powinna wchodzić nie tylko do kroku 1, ale też do kolejnych kroków, a ponadto trzeba te wejścia formalnie zapisać tak, by zapewnić sprawne i skuteczne korzystanie z odpowiednich dokumentów I. Dikmen et al. [2] zaproponowali koncepcję zarządzania ryzykiem przez uczenie się dla projektów budowlanych. Głównym zadaniem zaproponowanego przez nich narzędzia jest archiwizowanie informacji o rodzajach ryzyka i ich atrybutach, które pojawiły się w przeszłych projektach realizowanych przez firmy budowlane. Informacje te są gromadzone i aktualizowane w systemie w ciągu całego cyklu życia projektu. Autorzy narzędzia twierdzą, że dzięki analizie ryzyka z przeszłości łatwiej i skuteczniej jest oceniane ryzyko teraźniejsze, dzięki czemu decyzje dotyczące nowych projektów są bardziej racjonalne.

2. Podstawowe informacje o projektach europejskich

Przedmiotem dalszej części pracy będzie specjalna grupa projektów określanych jako projekty europejskie. Pod pojęciem projektów europejskich rozumiane są „projekty realizowane w ramach polityki i przy współdziałaniu Unii Europejskiej” [4]. W niniejszej pracy przyjmuje się, że projekty europejskie to projekty współfinansowane przez Unię Europejską, a realizowane przez polskie jednostki samorządu terytorialnego, zwłaszcza gminy.

Po wstąpieniu Polski do Unii Europejskiej tematyka dotycząca projektów europejskich stała się bardzo popularna. Polska jest już po pierwszym okresie programowania, który miał miejsce w latach 2004-2006 oraz w trakcie okresu programowania na lata 2007-2013. Niemniej jednak, pomimo zdobycia już pewnego doświadczenia w tym zakresie, dalej można zidentyfikować wiele błędów popełnianych przy realizacji i rozliczaniu projektów europejskich przez polskie gminy, wśród których przede wszystkim można wymienić: nieprawidłowe przygotowanie wniosku unijnego, nieprawidłowe przeprowadzenie procedury przetargowej, nieprzestrzeganie zasad kwalifikowalności poszczególnych kategorii wydatków, niedochowanie wymogów dotyczących informacji i promocji projektu, brak prawidłowej archiwizacji dokumentacji i w rezultacie niemożność udokumentowania poniesionych kosztów [3]. Konsekwencje wymienionych błędów mogą być ogromne i spowodować, że zwrot wydatków poniesionych przez gminę zostanie zakwestionowany przez jedną z licznych instytucji kontrolujących prawidłowość pozyskania i wykorzystania funduszy unijnych, czego konsekwencją będą opóźnienia w otrzymaniu środków unijnych, zmniejszenie ich kwoty, a nawet odmowa ich przyznania bądź konieczność zwrotu. Można więc stwierdzić, że projekty europejskie wyróżniają

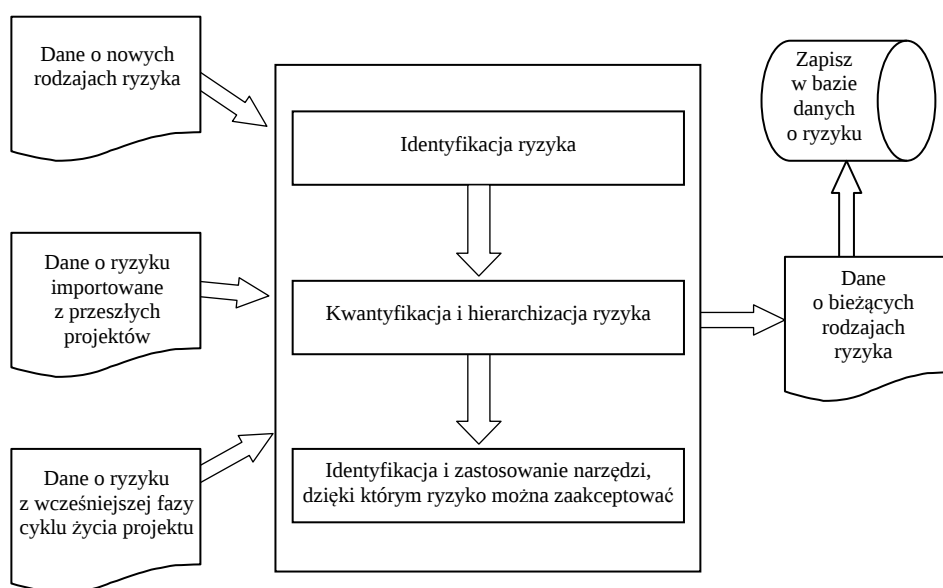
się wysokim stopniem ryzyka, które może być dotkliwe w skutkach, zwłaszcza z punktu widzenia społeczeństwa. W związku z tym konieczne jest skuteczne i efektywne zarządzanie tego typu projektami i ryzykiem z nimi związanym. Pomimo tego, że w literaturze z zakresu zarządzania projektami często jako jedną z głównych cech projektu wyróżnia się jego unikalność, to wśród projektów europejskich realizowanych w gminach są grupy projektów do siebie podobnych. Podobieństwo wynika głównie z programu, z jakiego są współfinansowane (np. program INTERREG), z celu działania, jakiego dotyczą (np. odbudowy dróg), a także ich dziedziny (np. budownictwo). Ważne jest więc, aby gminy wykorzystywały doświadczenia zdobyte w przeszłości przy podobnych typach projektów i korygowały swoje działania w zależności od tego, czy wcześniejsze działanie przyniosło sukces czy porażkę. Bez uczenia się zarządzanie będzie znacznie bardziej utrudnione. Stąd w kolejnym rozdziale zostanie opisana propozycja systemu zarządzania ryzykiem projektów europejskich przez uczenie się, który zrealizowany byłby za pomocą odpowiednich narzędzi informatycznych (MS Access).

3. Propozycja systemu zarządzania projektami europejskimi przez uczenie się

Proponowany system zarządzania ryzykiem projektów europejskich obejmowałby pięć modułów: projekty zainicjowane, projekty planowane, projekty wnioskowane, projekty realizowane i projekty zakończone. Każdy z wymienionych modułów stanowi fazę cyklu życia projektu europejskiego w gminie. Najpierw jest więc faza zainicjowania projektu, w której inicjator projektu przedstawia pomysł na projekt i ogólny sposób jego realizacji wraz ze wskazaniem korzyści z niego wynikających. W tej fazie należy przede wszystkim określić w gminie osoby, które będą odpowiedzialne za gromadzenie i aktualizowanie danych na temat nowego projektu w systemie oraz wprowadzić do systemu podstawowe informacje o nowym projekcie i jego udziałowcach, a także przypisać go do odpowiedniego typu i podtypu projektów europejskich realizowanych w gminie. Jeżeli pomysł na projekt zostanie zaakceptowany na posiedzeniu Rady Gminy, to przystępuje się do fazy planowania projektu, czyli stworzenia szczegółowego planu jego realizacji, który ze względu na wymagania formalne Unii Europejskiej jest dość sztywnym planem i niemalże niemożliwym do modyfikacji. Kolejną fazą cyklu życia projektu europejskiego w gminie jest faza przygotowania wniosku unijnego. Po zaakceptowaniu planu projektu oraz uzyskaniu zgody na dotację unijną przystępuje się do faktycznego wykonania projektu i w końcu ostatnią fazą cyklu życia projektu europejskiego w gminie jest jego zakończenie wraz ze złożeniem sprawozdania z przebiegu

realizacji projektu do instytucji UE. Każda z wymienionych pięciu faz w każdym momencie może zostać przerwana i skończyć się niepowodzeniem, a więc stać się ostatnią fazą cyklu życia projektu europejskiego. Użytkownik rozpoczynający pracę z systemem może zainicjować nowy projekt lub wyświetlić listę wszystkich projektów i wybrać jeden z nich, będący już w dalszej fazie zaawansowania.

W proponowanym systemie zarządzanie ryzykiem odbywa się w każdej fazie cyklu życia projektu europejskiego, w oparciu o różne źródła informacji (rys. 2).



Rys. 2. Schemat działania proponowanego systemu zarządzania ryzykiem projektów europejskich

Źródło: Opracowanie własne na podstawie [1].

Prawie we wszystkich fazach cyklu życia projektu europejskiego zarządzanie ryzykiem dla danego projektu odbywa się na podstawie zaktualizowanych danych o ryzyku z wcześniejszej fazy oraz danych o ryzyku importowanym z podobnych projektów realizowanych w przeszłości przez własną gminę lub inną gminę uczestniczącą w systemie. Poza ryzykiem zaproponowanym przez system i importowanym z przeszłych projektów użytkownik w każdym module może wprowadzać i opisywać nowe rodzaje ryzyka. Nieco inaczej jest w przypadku projektu wnioskowanego, gdzie zarządzanie ryzykiem odbywa się na podstawie danych o ryzyku, które pojawiło się przy wypełnianiu podobnych wniosków w przeszłości. Dodatkowo system sam proponuje rodzaje ryzyka i ich opis, które są domyślnie przyporządkowane dla typu programu unijnego, do którego można zaliczyć analizowany projekt.

Reasumując, głównym zadaniem proponowanego systemu będzie gromadzenie, przechowywanie, aktualizowanie i dzielenie się informacjami o ryzyku występującym w projektach europejskich, które są realizowane przez polskie gminy w coraz większej liczbie i angażujących coraz większe zasoby.

4. Problemy występujące w praktyce polskich gmin

Na podstawie rozmów z pracownikami i władzami wybranych gmin w Polsce oraz analizy dokumentacji dotyczącej konkretnych projektów zidentyfikowano główne problemy pojawiające się w zarządzaniu ryzykiem projektów europejskich w polskich gminach.

Jednym z najpoważniejszych problemów polskich gmin przy realizacji projektów europejskich jest zbyt pobieżne traktowanie ryzyka takich projektu. Pracownicy gmin, sporządzający wniosek o dotację, zwykle koncentrują się wyłącznie na ryzyku sporządzenia nieprawidłowego wniosku unijnego. Chcąc za wszelką cenę otrzymać dotację, bagatelizują ryzyko niepowodzenia całego projektu. Jedynym elementem zarządzania ryzykiem we wnioskowanych projektach europejskich jest Studium Wykonalności, które między innymi zawiera analizę opłacalności projektu oraz analizę ryzyka i wrażliwości. Niemniej jednak analiza ryzyka i wrażliwości przeprowadzona jest bardzo powierzchownie. Najpierw na podstawie analizy wrażliwości przyjmowane są trzy krytyczne czynniki mające największy wpływ na powodzenie lub porażkę projektu, a następnie zostają im przypisane oceny prawdopodobieństwa ich wystąpienia wyrażone słownie jako „wysokie”, „średnie” lub „niskie”. Ponadto, polskie gminy najczęściej nie są zaangażowane w analizę informacji zawartych w Studium Wykonalności, gdyż jest ono sporządzane przez firmy zewnętrzne zajmujące się taką działalnością. Takie bagatelizowanie zarządzania ryzykiem przez polskie gminy sprawia, że później w trakcie realizacji projektów pojawiają się problemy, których można by było uniknąć, gdyby wcześniej je przeanalizowano. Przykładowo, w jednej z polskich gmin, planując projekt budowy szkoły podstawowej, w trakcie sporządzania harmonogramu projektu nie uwzględniono ryzyka niezrealizowania wykonawcy. Okazało się, że gdy wniosek o dotację został rozpatrzony pozytywnie i przystąpiono do faktycznej realizacji projektu pojawił się problem ze znalezieniem firmy budowlanej, która podjęłaby się wykonania projektu na zasadach określonych przez gminę. Spowodowało to wydłużenie procedury przetargowej prawie o 1 rok, a tym samym wydłużenie całkowitego czasu realizacji projektu i utratę możliwości dofinansowania projektu ze środków unijnych. Kolejnym problemem jest fakt, że w polskich gminach brakuje właściwej dokumentacji i archiwizacji informacji na temat realizowanych projektów. Przykładowo, w jednej z gmin miała miejsce sytuacja, że osoba za-

rzządzają projektem europejskim w gminie z powodu choroby nie była w stanie kontynuować swojej pracy. Okazało się wówczas, że w gminie nie ma pracownika, który bez pomocy tej osoby potrafiłby kontynuować projekt, przez co gmina została narażona na stratę 7,5 mln złotych.

5. Potencjalne korzyści z zastosowania proponowanego systemu zarządzania ryzykiem

Główną korzyścią wynikającą z zastosowania proponowanego systemu zarządzania ryzykiem jest pomoc pracownikom polskich gmin w zarządzaniu ryzykiem projektów europejskich, a więc pomoc nie tylko w identyfikowaniu ryzyka, ale także w ograniczaniu bądź eliminowaniu poszczególnych rodzajów ryzyka pojawiającego się w różnych fazach cyklu życia takich projektów. Dzięki zastosowaniu systemu pracownicy mogliby przeanalizować ryzyka oraz ich skutki, które pojawiły w podobnych typach projektów w przeszłości i już na etapie ich wstępnego planowania wybrać projekty, dla których poziom całkowitego ryzyka jest najmniejszy. Wówczas tylko dla wybranych projektów byłby sporządzany wniosek unijny i ich szczegółowy plan. Ważną również korzyścią z zastosowania proponowanego systemu jest zarządzanie ryzykiem w fazie sporządzania wniosku unijnego. W polskich gminach nadal można wymienić wiele błędów popełnianych we wnioskach o dotację, których najczęstszą przyczyną jest niecelowe i wynikające z nieuwagi pominięcie pewnych istotnych informacji zawartych w obszernych instrukcjach. Błędnie wypełniony wniosek zostaje odrzucony lub zwrócony do poprawy, przez co pracownicy gmin muszą poświęcać swój czas na jego uzupełnienie zamiast koncentrować się na nowych wnioskach. Poza tym wśród korzyści wynikających ze stosowania proponowanego systemu można wymienić fakt, że w przypadku zmiany osób zarządzających danym projektem europejskim łatwiej jest przekazać obowiązki i informacje o projekcie nowym osobom.

Podsumowanie

W niniejszym opracowaniu zaproponowano koncepcję zarządzania ryzykiem projektów europejskich zakładającą uczenie się, korzystanie z doświadczeń własnych i obcych. Koncepcja ta wymaga dalszych konkretyzacji, w szczególności stworzenia jednolitego formatu pozwalającego z jednej strony archiwizować doświadczenia, z drugiej korzystać z nich, nie będąc narażonym na zasypanie ogromną ilością informacji nieprzejrzystych i być może nieprzydatnych dla danego projektu bądź dublujących się.

Literatura

1. Courtot H. (1998). La gestion des risques dans les projets. Economica, Paryż.
2. Dikmen I., Birgonul M.t., Anac C., Tah J.H.M. Aouad G. (2008). A Tool for Post-project Risk Assessment. Automation for Construction, 18.
3. Martini E., Bąk M. (2005). Dotacja i co dalej? Zarządzanie, sprawozdawczość, kontrola, promocja i ewaluacja projektów współfinansowanych z funduszy strukturalnych. Twigger, Warszawa.
4. Trocki M., Grucza B. (2007). Zarządzanie projektem europejskim. Polskie Wydawnictwo Ekonomiczne, Warszawa.

Krzysztof S. Targiel

Akademia Ekonomiczna w Katowicach

WYKORZYSTANIE OPCJI REALNYCH W SZACOWANIU RYZYKA PRZEKROCZENIA PRZEZ KOSZTY PROJEKTU USTALONEJ WARTOŚCI

Wstęp

Zarządzanie ryzykiem staje się jednym z najistotniejszych elementów zarządzania. Także w zarządzaniu projektami jest obecny ten element. Zbiór wiedzy z zakresu zarządzania projektami PMBOK [17] instytutu PMI (Project Management Institute) dziewięć obszarów wiedzy, wśród których jest także zarządzanie ryzykiem. Ryzyko rozumiane obecnie, w pewnym dystansie do oryginalnej definicji F.H. Knighta, jako zagrożenie niezrealizowania założonego celu. W ramach zarządzania ryzykiem rozróżniamy takie działania jak analiza, sterowanie i kontrola ryzyka. Występujące ryzyko najdobitniej jest widoczne na rynku giełdowym. Tutaj bardzo szybko wypracowano metody jego transferu poprzez wykorzystanie niesymetrycznych instrumentów pochodnych. Przykładem takiego instrumentu są opcje. Opcje wykorzystywane były do tego typu działań od momentu pierwszych notowań na Chicago Board Options Exchange (CBOE) w 1973 roku. Z tym rokiem wiąże się też najbardziej znany model wyceny opcji – model Blacka-Scholesa [2]. Praca Blacka i Scholesa dotyczyła opcji finansowych. Bardzo szybko zauważono, iż opcje mogą dotyczyć nie tylko finansów, lecz także tych aspektów życia, w których istnieje pewne prawo, z którego nie mamy obowiązku skorzystać. Ten typ nazwano „opcjami realnymi”, a nazwa wiąże się z nazwiskiem Myersa [15]. Sytuacje, w których możemy z pewnego prawa skorzystać, często występują w trakcie zarządzania projektami. Zaczęto wykorzystywać opcje realne do zarządzania w nich ryzykiem. Celem pracy jest przedstawienie możliwości wykorzystania opcji realnych, a w szczególności metod ich wyceny do szacowania ryzyka przekroczenia przez koszty projektu ustalonej w kontrakcie wartości.

1. Opcje realne w zarządzaniu projektami

Myers jako pierwszy zauważył nieadekwatność oceny projektów przy pomocy zdyskontowanych przepływów kapitałowych [16]. Metody te znane pod nazwą DCF (Discounted Cash Flow Analysis) lub wykorzystujące bieżącą wartość NPV (Net Present Value), nie uwzględniają przyszłych zmian, także pozytywnych, w analizowanych przepływach kapitałowych. Reprezentują one podejście statyczne. Powody te doprowadziły Myersa do zaproponowania dynamicznie dostosowywanej wartości bieżącej APV (Adjusted Present Value). Wielkość ta uwzględniała nie tylko przewidywane przepływy kapitałowe, ale także możliwe zmiany wiodące do wzrostu wartości projektu. Były to możliwości do spożytkowania. Myers dla tych możliwości zaproponował sformułowanie „opcje realne” [15].

Trigeorgis [18] przedstawia wiele typów opcji realnych. Między innymi są to:

1. Opcja rezygnacji (option to abandon) – gdy warunki rynkowe sprawią iż projekt staje się nieopłacalny, zarząd może podjąć decyzję o rezygnacji z jego kontynuowania. Powyższą możliwość można traktować jako amerykańską opcję sprzedaży.

2. Opcja zmiany zakresu działania – ma dwa podtypy opcję rozszerzenia (option to expand), gdy warunki rynkowe są korzystniejsze niż zakładano, zarząd projektu może podjąć decyzję o zwiększeniu nakładów i dzięki temu zwiększeniu skali projektu, lecz dzięki temu zwiększeniu przyszłych profitów. Możliwość ta może być traktowana jako amerykańska opcja kupna. Przeciwnieństwem tej sytuacji jest opcja redukcji lub zmniejszenia skali działania (scope down option). Jest to amerykańska opcja sprzedaży.

3. Opcja rozwoju (growth option) możliwość wykorzystania aktywów do rozwoju nowoczesnych technologii, które w przyszłości mogą dać znaczące profity.

4. Opcja wydłużenia – (option to extend) możliwość wydłużenia życia produktu w korzystnych warunkach rynkowych (amerykańska opcja kupna).

5. Opcja skrócenia – (option to shorten) możliwość skrócenia życia produktu w niekorzystnych warunkach rynkowych (amerykańska opcja sprzedaży).

6. Opcja zamiany – (option to switch) możliwość innego wykorzystania aktywów zaangażowanych w projekt. Te same aktywa można wykorzystać do wytworzenia innych produktów (product flexibility). Ten podtyp opcji jest nazywany option to switch output. Alternatywnie, do wytworzenia tego samego produktu można wykorzystać inne aktywa (process flexibility). Ten podtyp opcji jest nazywany option to switch input. Tego typu opcje mogą być wyceniane jako amerykańskie opcje sprzedaży.

7. Opcja rozpoczęcia projektu (option to defer) – jest to możliwość opóźnienia rozpoczęcia projektu do momentu uzyskania nowych informacji skutkujących lepszą ceną opłacalności. Sytuacja taka może być wyceniana jako amerykańska opcja kupna.

Powyższe możliwości tworzą dla firmy nową wartość. Oprócz realnych aktywów, posiada ona także realne możliwości (opcje). Do wyceny tych dodatkowych realnych możliwości wykorzystywane są najczęściej drzewa dwumianowe [18], ale także modyfikacje wzoru Blacka-Scholesa [1; 7]. Pojawiły się także propozycje wyceny opcji realnych w oparciu metodologię zbiorów rozmytych [3; 4; 5].

2. Opcje realne a zarządzanie ryzykiem projektu

Podstawowe zastosowanie opcji realnych do wyceny projektów na etapie ich projektowania zostało uzupełnione o wykorzystanie tych instrumentów w zarządzaniu ryzykiem. Kumar [11] wykorzystuje opcje w zarządzaniu ryzykiem projektów informatycznych. Proponuje rozróżnienie sytuacji, które wymagają działania i sytuacji, w których można się przed skutkami zabezpieczyć. Do aktywnego zabezpieczenia się przed ryzykiem, co odpowiada postępowaniu znanemu w finansach pod nazwą hedging, wykorzystuje na różnych etapach życia projektu opcje: rozszerzenia, opóźnienia czy też rezygnacji. Wycena opcji jest wykonywana na podstawie formuły Margrabiego [12]. Benaroch i Kauffman [1] rozważają opcję opóźnienia projektu rozwoju informatycznej sieci bankowej. Do wyceny wykorzystują model Blacka-Scholesa. Wu, Ong i Hsu [19] rozważają projekty implementacji systemów ERP z perspektywy opcyjnej. Sytuacja taka jest modelowana jako opcja złożona. Do jej wyceny autorzy wykorzystują drzewa dwumianowe. Datar i Mathews [7] opracowali intuicyjną metodę wyceny opcji realnych. Metoda, która jak twierdzą autorzy jest równoważna metodzie Blacka-Scholesa, jest wykorzystywana w koncernie Boeinga do analizy ryzyka zaawansowanych technologicznie projektów. Praca Meinla i Neumanna [13] dotyczy wykorzystania opcji realnych do zabezpieczania dostępności mocy obliczeniowych i ograniczania ryzyka z tym związanego. Powyższy krótki przegląd zastosowań opcji realnych w zarządzaniu ryzykiem dotyczył projektów informatycznych, ponieważ takich projektów będzie dotyczyła także proponowana w dalszej części pracy metoda.

3. Model oceny ryzyka w projekcie

Przedstawiamy model służący ocenie prawdopodobieństwa zakończenia się projektu fiaskiem z powodu przekroczenia przez koszty wartości kontraktu. Jest to model, którego inspiracją była metoda zaproponowana przez Mertona [14], opisana w pracach [8; 9], służąca do oceny prawdopodobieństwa upadku firmy.

Rozważmy projekt informatyczny, w którym zakontraktowano stworzenie oprogramowania. Umowa ma ustaloną cenę, po której produkt (oprogramowanie) będzie sprzedane. Projekt trwa dwa lata. Zakładamy, że po pierwszym roku można ocenić parametry dynamiki zmian kosztów ponoszonych w stosunku do zaplanowanych.

Zajmujemy się projektem informatycznym ze względu na jego specyfikę, szczególnie przydatną w proponowanej metodzie. W projektach informatycznych, w których chodzi o utworzenie nowego oprogramowania, generowane koszty związane są z pracą programistów, analityków, testerów. Są to w głównej mierze koszty pracy. Nie występują koszty zakupu materiałów, urządzeń itd. Stąd rozważane koszty są trudne do dokładnego oszacowania. Ich dynamikę będziemy starali się opisać za pomocą geometrycznego procesu Wienera.

Przyjęto następujące oznaczenia:

- U – ustalona cena kontraktu,
- t – czas do zakończenia projektu (jako część roku),
- T – moment zakończenia projektu,
- Z_T – zysk z przedsięwzięcia po zakończeniu projektu,
- P_t – zaplanowane skumulowane koszty w kolejnych okresach,
- C_t – skumulowane koszty poniesione w kolejnych okresach,
- K_t – koszty dodatkowe w kolejnych okresach,
- V_t – wartość projektu w kolejnych okresach.

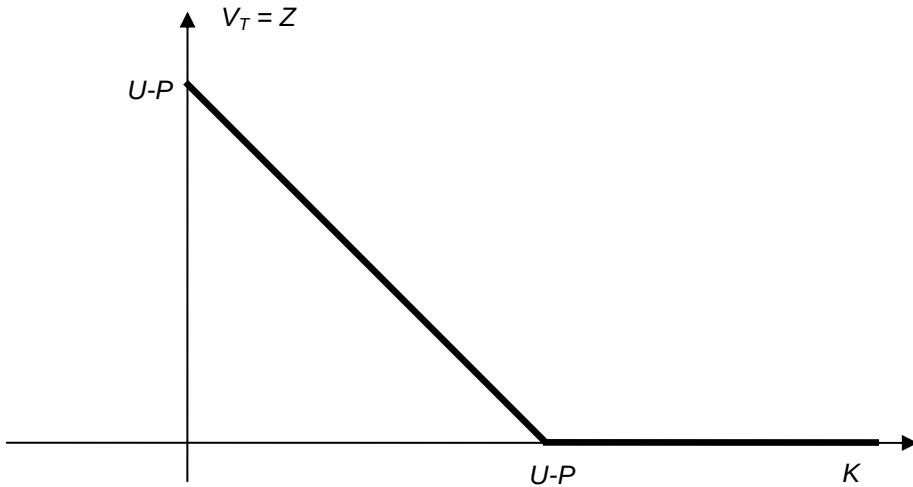
Gdy koszty przekroczą ustaloną cenę kontraktu, zysk z przedsięwzięcia będzie zerowy, co można zapisać w postaci

$$Z_T = \max\{(U - (P_T + K_T)), 0\} \quad (1)$$

lub w równoważnej postaci

$$Z_T = \max\{(U - P_T) - K_T, 0\} \quad (2)$$

Z punktu widzenia menadżera projektu powyższą sytuację można uważać za realną opcję sprzedaży o wypłacie przedstawionej na rys. 1.



Rys. 1. Zyski z projektu po jego zakończeniu jako wypłata z opcji sprzedaży

Dodatkowe niezaplanowane koszty (K_t), obliczane są jako różnica pomiędzy logarytmem kosztów poniesionych a logarytmem kosztów zaplanowanych

$$\ln K_t = \ln C_t - \ln P_t \quad (3)$$

Wtedy przyjmujemy, że dodatkowe koszty zmieniają się zgodnie z następującym równaniem dynamiki

$$dK_t = \mu \cdot K_t dt + \sigma_K \cdot K_t dW \quad (4)$$

W tej sytuacji wartość projektu można obliczyć z modelu Blacka-Scholesa dla europejskiej opcji sprzedaży, to znaczy zgodnie z poniższymi wzorami [6]

$$V_t = -K_t \cdot \Phi(-d_1) + (U - P_T) \cdot e^{-rt} \cdot \Phi(-d_2) \quad (5)$$

$$d_1 = [\ln(K_t / (U - P_T)) + (r + \sigma_K^2 / 2) \cdot t] / [\sigma_K \cdot (t)^{0,5}] \quad (6)$$

$$d_2 = d_1 - \sigma_K \cdot (t)^{0,5} \quad (7)$$

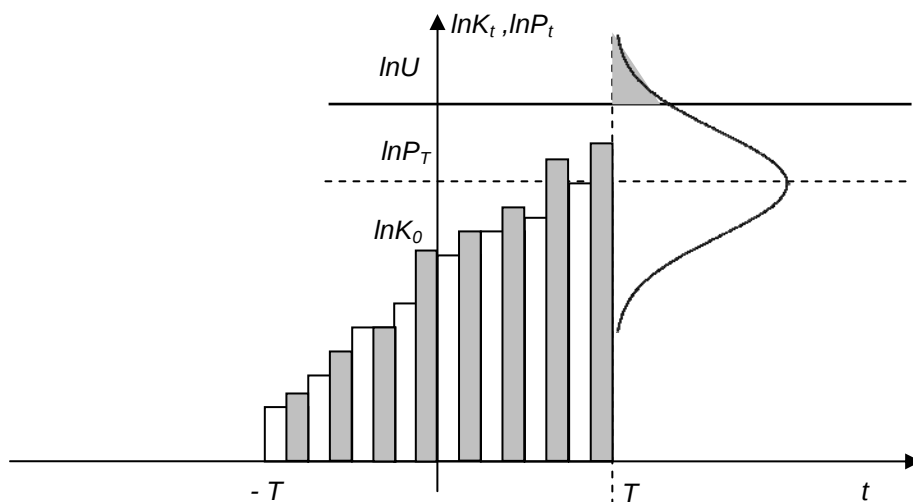
Konsekwencją przyjętego modelu dynamiki zmian wartości kosztów jest fakt, iż w momencie wykonania, logarytm wartości kosztów dodatkowych ($\ln K_T$) jako zmienna losowa, ma rozkład normalny

$$\ln K_T \sim N(\ln K_0 + (\mu + \sigma_K^2 / 2) \cdot t ; \sigma_K \cdot (t)^{0,5}) \quad (8)$$

Stąd można obliczyć prawdopodobieństwo, iż koszty przekroczą po zakończeniu projektu ustaloną wartość kontraktu

$$P(\text{straty}) = P(\ln K_T \geq \ln(U - P_T)) = 1 - \Phi([\ln(U - P_T) - E(\ln K_T)] / \sigma(\ln K_T)) \quad (9)$$

Przykładowy proces kosztów, obserwowanych w okresach kwartalnych, przedstawiono na rys. 2.



Rys. 2. Proces zmian kosztów projektu

Zaplanowany poziom kosztów jest na rys. 2 oznaczony niewypełnionymi słupkami. Po jednym roku trwania projektu koszty powinny osiągnąć poziom K_0 . Słupki wypełnione przedstawiają rzeczywiste poniesione koszty. Postawione pytanie jest o prawdopodobieństwo z jakim rzeczywiście poniesione koszty mogą one przekroczyć poziomu wartości kontraktu (U), tzn. znaleźć się w szarym obszarze. Wykres przedstawiono w skali logarytmicznej, ponieważ łatwiej jest wtedy narysować funkcję gęstości rozkładu prawdopodobieństwa. W tej skali jest to bowiem w krzywa gaussowska. Na przedstawionym rysunku pokazano realizacje procesu kosztów jedynie dla okresów kwartalnych, by nie zaciemniać rysunku. W rzeczywistych obliczeniach należy jednak brać pod uwagę co najmniej miesięczne przyrosty kosztów, jeżeli nie częstsze.

Wykorzystując dystrybucję rozkładu normalnego Φ prawdopodobieństwo poniesienia straty wynosi

$$P(\text{straty}) = 1 - \Phi([\ln(U - P_T) - (\ln K_0 + (\mu + \sigma_K^2/2) \cdot t)] / (\sigma_K \cdot t)^{0.5}) \quad (10)$$

Aby obliczyć prawdopodobieństwo straty, konieczna jest znajomość parametrów geometrycznego procesu Wienera opisującego dynamikę zmian wartości kosztów, parametru dryfu μ oraz parametru zmienności kosztów σ_K .

Wykorzystać można model trendu liniowego przechodzącego przez początek układu współrzędnych. Współczynnik kierunkowy modelu, można traktować jako parametr μ modelu dynamiki. Odchylenia od modelu trendu, można traktować jako realizację składnika losowego $\sigma_K dW$ w modelu dynamiki. Odchylenie standardowe reszt można traktować jako estymator zmienności kosztów σ_K .

4. Przykład numeryczny

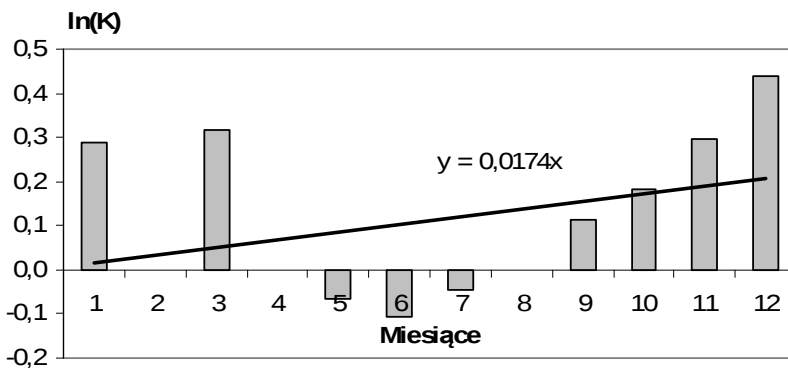
Obliczenia ilustrujące metodę, przeprowadzono na hipotetycznych danych przedstawiających przebieg zaplanowanych skumulowanych kosztów (P_t) i poniesionych skumulowanych kosztów (C_t) w kolejnych miesiącach pierwszego roku trwania projektu. Dane te przedstawiono w tabeli 1. Umowa zawarta na wykonanie oprogramowania zawiera kwotę $U = 10$ mln PLN płatną po przyjęciu produktu przez zamawiającego. Zaplanowano, że koszty wytworzenia (C_T) będą na poziomie 8 mln PLN.

Tabela 1

Dane do obliczeń

Miesiące	1	2	3	4	5	6	7	8	9	10	11	12
P_t	0,3	0,5	0,8	1,2	1,6	2	2,2	2,3	2,5	3	3,5	4
C_t	0,4	0,5	1,1	1,2	1,5	1,8	2,1	2,3	2,8	3,6	4,7	6,2
$\ln P_t$	-1,2	-0,7	-0,2	0,2	0,5	0,7	0,8	0,8	0,9	1,1	1,3	1,4
$\ln C_t$	-0,9	-0,7	0,1	0,2	0,4	0,6	0,7	0,8	1,0	1,3	1,5	1,8
$\ln K_t$	0,3	0,0	0,3	0,0	-0,1	-0,1	0,0	0,0	0,1	0,2	0,3	0,4

Wykres różnic pomiędzy logarytmami kosztów poniesionych (C_t) a logarytmami kosztów zaplanowanych (P_t) przedstawia rys. 3.



Rys. 3. Różnice pomiędzy kosztami poniesionymi a zaplanowanymi

Dla przedstawionego szeregu czasowego wyestymowano metodą najmniejszych kwadratów model trendu liniowego, przechodzącego przez początek układu współrzędnych. Wyznacza on średni poziom wzrostu kosztów ponad zaplanowane. Jego wykres jest także umieszczony na rys. 3. Współczynnik kierunkowy modelu potraktowano jako parametr μ modelu dynamiki. Odchylenia od modelu trendu to realizacja składnika losowego $\sigma_K dW$ w modelu dynamiki. Odchylenie standardowe reszt uznano za estymator zmienności kosztów σ_K .

Po użyciu metody najmniejszych kwadratów uzyskano

$$\begin{aligned}\mu &= 0,0174 \\ \sigma_K &= 0,1732\end{aligned}$$

Po wykorzystaniu wzoru (10), uzyskujemy prawdopodobieństwo faktu, że koszty przekroczą za rok wartość kontraktu

$$P(\text{straty}) = 1 - \Phi(1,2847)$$

Stąd prawdopodobieństwo to wynosi

$$P(\text{straty}) = 0,0994$$

Ryzyko przekroczenia przez koszty poziomu przychodów wynosi około 10%. W obliczeniach przyjęto $K_0 = 2,2$ mln PLN.

Podsumowanie

W pracy przedstawiono metodę oceny ryzyka zakończenia się przedsięwzięcia porażką na skutek przekroczenia przez koszty poziomu zagwarantowanych przychodów z tytułu zawartej umowy. Powyższą sytuację opisano jako opcję realną, w której na skutek niekorzystnego przebiegu zdarzeń polegającego na znacznym wzroście kosztów, przedsięwzięcie traci swoją wartość. Wykorzystując znane z finansów, metody wyceny opcji, uzyskano oszacowanie ryzyka zakończenia się projektu niepowodzeniem. Jest to wstępna propozycja tej metody. W szczególności wymaga potwierdzenia, np. na podstawie badań empirycznych, słuszność wyboru modelu dynamiki zmian ponoszonych kosztów. Dopracowane powinny zostać metody oceny parametrów procesu. Wydaje się jednak, iż metoda może być skutecznie wykorzystywana do proaktywnego zarządzania ryzykiem w projekcie, zwłaszcza informatycznym, w którym koszty pracy mogą rosnąć lawinowo i przyczynić się do niepowodzenia projektu. Na ich potencjalnie niebezpieczny wzrost można reagować już w momencie, gdy pojawia się taka możliwość, czyli realna opcja.

Literatura

1. Benaroch M., Kauffman R.J. (2000). Justifying electronic Banking Network Expansion Using Real Options Analysis. *MIS Q*, 24, 2, 197-225.
2. Black F., Scholes M. (1973). The Pricing of Option and. Corporate Liabilities. *Journal of Political Economy*, 81, 3, 637-659
3. Carlsson C., Fuller R. (2003). A Fuzzy Approach to real Option Valuation. *Fuzzy Sets and Systems*, 139, 297-312.
4. Carlsson C., Fuller R., Heikkila M., Majlender P. (2007). A Fuzzy Approach to R&D project Portfolio Selection. *International Journal of Approximate Reasoning*, 44, 93-105.
5. Collan M.: Fuzzy Pay-Off Method For Real Option Valuation. <http://web.abo.fi/~mcollan/fuzzypayoff.html>
6. Daigler R.R. (1994). *Advanced Options Trading*. Probus Publishing Company, Chicago.
7. Datar V.T., Mathews S.H.: European Real Options: An Intuitive Algorithm for the Black-Scholes Formula. SSRN eLibrary. http://papers.ssrn.com/sol3/papers.cfm?abstract_id=560982
8. Gątarek D., Maksymiuk R. i inni. (2001). *Nowoczesne metody zarządzania ryzykiem finansowym*. WIG-Press, Warszawa.
9. Jajuga K. (2007). *Zarządzanie ryzykiem*. Wydawnictwo Naukowe PWN, Warszawa.
10. Jakubowski J. i i inni. (2003). *Matematyka finansowa. Instrumenty pochodne*. WNT, Warszawa.
11. Kumar R.L. (2002). Managing Risks in IT projects: An Options Perspective. *Information & Management*, 40, 63-74.
12. Margrabe W. (1978). The Value of an Option to Exchange One Asset for Another. *Journal of Finance*, 33, 1, 177-186.
13. Meinel Th., Neumann D. (2009). A Real Options Model for Risk Hedging in Grid Computing Scenarios. W: *Proceedings of the 42nd Hawaii International Conference on System Sciences*.
14. Merton R.C. (1974). On the Pricing of Corporate Debt: the Risk Structure of Interest Rates. *Journal of Finance*, 29, 449-470.
15. Myers S.C. (1974). Interactions of Corporate Financing and Investment Decisions-Implications for Capital Budgeting. *The Journal of Finance*, 29, 1, 1-25.
16. Myers S.C. (1977). Determinants of Corporate Borrowing. *Journal of Financial Economics*, 5, 2, 147-175.

17. Project Management Institute (2004). A Guide to the Project Management Book of Knowledge (PMBOK). Project Management Institute, Newtown Square, PA.
18. Trigeorgis L. (1993). Real Options and Interactions with Financial Flexibility. *Financial Management*, 22, 3, 202-224.
19. Wu L-Ch, Ong Ch-Sh, Hsu Y-W. (2008). Active ERP Implementation Management: A Real Options Perspective. *The Journal of Systems and Software*, 81, 1039-1050.

Wojciech Walczak

Politechnika Wrocławska

ZARZĄDZANIE RYZYKIEM W ZWINNYCH METODYKACH ZARZĄDZANIA PROJEKTAMI

Wprowadzenie

W metodykach zwinnych występuje dość specyficzne podejście do procesów, w tym do zarządzania ryzykiem. Dzieje się tak, ponieważ metodyki te przystosowane są do kontekstu projektu, w którym stale zachodzą zmiany, a przyszłe okoliczności są nieprzewidywalne. Zatem nie jest znane środowisko, w którym miałyby zachodzić proces zarządzania ryzykiem, czyli nie jest możliwe zdefiniowanie procesu, który byłby do niego dopasowany. Jednoczesne kładzenie przez te metodyki nacisku na generowanie wartości dodanej sprawiło, że wszelkie prace niezwiązane bezpośrednio z budowaniem funkcjonalności tworzonego systemu są traktowane niechętnie. Wymienione czynniki nie sprzyjają istnieniu pełnego procesu zarządzania ryzykiem w metodykach zwinnych.

Mówiąc o zarządzaniu ryzykiem nie należy zapominać o rozgraniczeniu między ryzykiem a niepewnością. O ryzyku mówimy w kontekście niepewnych zdarzeń, wpływających pozytywnie lub negatywnie na przebieg projektu i których prawdopodobieństwo można oszacować. W przypadku niepewności prawdopodobieństwo zdarzenia jest zupełnie nieznane. Z pewnością jest bardzo wiele niepewności w projektach, dla których stworzone zostały metodyki zwinne. Jednakże, nawet w bardzo dynamicznym środowisku, zawsze obecna jest pewna pula zagrożeń, których przewidzenie jest możliwe, choćby z niewielkim wyprzedzeniem. Możliwe byłoby wcześniejsze opracowanie dla nich zestawu symptomów ułatwiających diagnozę oraz wprowadzenie działań zwalczających ryzyko, jeszcze zanim się ono zmaterializuje. Przykładem takich zagrożeń mogą być wszystkie te niezwiązane z otoczeniem biznesowym, które zrodziło projekt, lecz z organizacją, w której projekt jest osadzony. Czy istnienie ogromnej ilości niepewności w projekcie zwalnia z obowiązku prób identyfikacji ryzyka, choćby niewielkiego, lecz niewykluczone, że bardzo prawdopodobnego i katastro-

falnego w skutkach? Okazuje się, że nie i z tej przyczyny została podjęta próba uzupełnienia metodyk zwinnych o praktyki, które będą wprowadzały do nich zarządzanie ryzykiem.

Celem niniejszego opracowania jest skonstruowanie procesu zarządzania ryzykiem, który będzie mógł być stosowany wraz z metodykami zwinnymi zarządzania projektami, nie pozbawiając ich zalet i cech charakterystycznych.

1. Metodyki zwinne

Za początek zwinnego zarządzania projektami uznaje się powstanie Manifestu Zwinnego Wytwarzania Oprogramowania <http://agilemanifesto.org/> w lutym 2001, choć niektóre metodyki, obecnie określane mianem zwinnych, istniały już wcześniej. Metodyki te traktują zmiany w projekcie jako nieusuwalny i niezbędny element towarzyszący projektowi, a nawet dostrzegają pozytywną stronę niektórych rodzajów zmian. Skupiają się one na generowaniu wartości dodanej, skróceniu czasu oddania produktu, częstych interakcjach z klientem, pracy w krótkich iteracjach, samoorganizujących się zespołach, usprawnieniu komunikacji wewnątrz zespołu oraz eliminacji niepotrzebnej dokumentacji projektowej.

Niżej zostały zwięźle przedstawione dwie najbardziej powszechne metodyki zwinne: SCRUM i eXtreme Programming (XP). Na podstawie ich cech wspólnych zbudowany został model zwinnego procesu wytwarzania oprogramowania.

1.1. SCRUM

Metodyka SCRUM [3] skupia się na organizacji zespołów pracujących w projektach o ciągle zmieniającym się otoczeniu. Przed właściwym wytwarzaniem oprogramowania, przeprowadzana jest faza rozpoczęcia, której rezultatami są: wstępna definicja zakresu projektu, wstępny plan edycji produktu oraz architektura wysokiego poziomu.

Zakres jest opisany przez rejestr wymagań (Product Backlog), który oprócz samych wymagań, zwykle spisanych jako historyjki użytkownika (user stories), zawiera również ich priorytety nadane przez klienta. Rejestr wymagań nie jest dokumentem zamkniętym i w trakcie projektu jest często aktualizowany. Plan edycji produktu również może ewoluować w późniejszych etapach projektu.

Faza wytwarzania składa się z szeregu stałej długości iteracji (sprintów), zwykle czterotygodniowych. Każda iteracja rozpoczyna się wyborem wymagań z rejestru wymagań – zadania niezbędne do ich dostarczenia zostają wpisane

do rejestru iteracji (Sprint Backlog). Rejestr wymagań, rejestr cyklu oraz wykres szybkości realizacji prac w iteracji, stanowią jedyną wymaganą przez metodykę dokumentację projektową. W każdej iteracji przeprowadzone są wszystkie kroki, niezbędne do dostarczenia wybranej funkcjonalności, czyli wykonana jest analiza wymagań, przygotowane są szczegóły projektu oprogramowania, wprowadzone niezbędne zmiany do już istniejących elementów produktu a nowa funkcjonalność jest implementowana i testowana. Ważne jest, że wymagania nie mogą być modyfikowane w trakcie trwania iteracji, jeśli opisana przez nie funkcjonalność jest aktualnie budowana. Kończąc iterację, członkowie zespołu zbierają się na spotkaniu retrospektywnym, gdzie omawiają przebieg ostatniej iteracji oraz rozważają udoskonalenia procesu wytwarzania oprogramowania. Produktem każdej iteracji jest oprogramowanie, które potencjalnie może być dostarczone do klienta. Toteż wraz z końcem projektu, jako ostateczny produkt dostarczony jest po prostu produkt ostatniej iteracji.

1.2. eXtreame Programming

Metodyka eXtreame Programming (XP) pomaga w osiągnięciu celu nadzrzednego, jakim jest sukces projektu, przedstawiając zestaw praktyk, które powinny być stosowane w trakcie wytwarzania oprogramowania. Wskazanych jest trzynaście praktyk podstawowych i jedenaście praktyk uzupełniających [4]. W XP można wyróżnić pięć faz: eksploracji, planowania, iteracji, produktywacji, utrzymania oraz śmierci.

Faza eksploracji obejmuje okres od kilku tygodni do kilku miesięcy i skupia się na dostarczeniu wymagań przez użytkownika, które zwykle spisywane są w postaci historyjek użytkownika. Równolegle zespół projektowy zaznajamia się z technologią, narzędziami i praktykami, którym mają się podporządkować. W fazie planowania, trwającej kilka dni, zespół projektowy wraz z klientem oceniają możliwości produkcyjne. Programiści oceniają nakłady pracy potrzebne do realizacji wymagań użytkownika, a osoba pełniąca rolę przywódczą w zespole zarysowuje grafik obejmujący nie więcej niż dwa następne miesiące.

Faza iteracji obejmuje kilka cykli tworzenia oprogramowania, które nie powinny być dłuższe niż kilka tygodni, a zalecane są cykle jednotygodniowe. Każda iteracja kończy się testami funkcjonalnymi. Proponowane są również cykle kwartalne, będące szeregiem tychże jednotygodniowych iteracji i w których rozpatruje się cele projektu i problemy na wyższym poziomie uogólnienia.

Dodatkowe testy (wydajnościowe i akceptacyjne) przeprowadzane są w fazie produktywacji, kończącej się rozmieszczeniem systemu w środowisku klienta. Na tym etapie nadal mogą być wprowadzone nowe wymagania do

produktu. Faza utrzymania również jest szeregiem iteracji, skupiających się na wprowadzeniu do produktu zmian i nowych funkcjonalności. Faza śmierci rozpoczyna się z chwilą, gdy rozwijanie systemu przestaje być opłacalne dla klienta, gdy wyczerpują się u klienta wymagania poprawy, udoskonalenia lub dostosowania produktu. W fazie tej tworzone są wszystkie niezbędne dokumenty a oprogramowanie jest ostatecznie rozmieszczane u klienta.

1.3. Model procesu zwinnego wytwarzania oprogramowania



Rys. 1. Model procesu wytwarzania oprogramowania według metodyki zwinnej

Stworzony model procesu wytwarzania oprogramowania według metodyki zwinnej widoczny jest na rys. 1. Zawiera on elementy najczęstsze i najbardziej charakterystyczne dla metodyk zwinnych.

2. Zagrożenia adresowane przez metodyki zwinne

Częste dostawy kolejnych fragmentów oprogramowania oraz stały kontakt z klientem pozwalają w znacznym stopniu wyeliminować ryzyko wynikające ze stale zmieniającego się otoczenia projektu. Zmiany takie mogą spowodować znaczne opóźnienia lub koszty w przypadku tradycyjnych metodyk, gdyż nie jest możliwa dostatecznie szybka reakcja na nie. W metodykach zwinnych procesy nie mają takich ograniczeń.

Metodyki zwinne niwelują również ryzyko nieprawidłowych wymagań wobec produktu lub ich dezaktualizacji w trakcie trwania projektu. Dostawca oprogramowania ma stale aktualizowaną listę wymagań, których priorytety są na bieżąco określane przez klienta i to właśnie z tej listy, a nie z zamkniętej specyfikacji stworzonej na początku projektu, pobierane są wymagania podczas planowania każdej iteracji. Dzięki temu implementowane funkcjonalności odpowiadają aktualnym potrzebom i oczekiwaniom klienta, a nie tym z chwili podpisania kontraktu. Jak łatwo zauważyć, wyeliminowane zostało również ryzyko implementacji funkcjonalności niepotrzebnych.

Ryzyko braku zdecydowania klienta lub braku umiejętności wyartykułowania swoich potrzeb i oczekiwań jest pomniejszane poprzez częste dostawy kolejnych funkcjonalności. Prace nie są wstrzymywane na czas tworzenia dobrze wyspecyfikowanych wymagań, w pełni zgodnych z faktycznymi oczekiwaniami. Klient może doprecyzowywać swoje wymagania w trakcie wytwarzania oprogramowania, a wraz z końcem każdej iteracji ma okazję sprawdzić, czy gotowy fragment produktu mu odpowiada oraz zgłosić swoje ewentualne uwagi.

Oprócz tego, szybszy czas wejścia produktu na rynek, nawet niekompletnego, zmniejsza ryzyko biznesowe utracenia udziałów w rynku i pomaga uzyskać przewagę nad konkurencją. Natomiast poprzez położenie nacisku na bezpośredni kontakt międzyludzki, metodyki zwinne starają się zapobiec ryzyku nieporozumień w zespole oraz błędnych interpretacji wszelkich treści przez członków zespołu.

3. Metodyki zwinne a zarządzanie ryzykiem

Dość popularne są poglądy, zwłaszcza wśród praktyków zarządzania zwinnego, że jeśli stosowane są metodyki zwinne to zarządzanie ryzykiem jest zbędne. Nie są one zawieszone w próżni – w literaturze możemy znaleźć opinie, że krótkie iteracje i oparcie się na decyzjach klienta przy wyborze funkcjonalności, która ma być zaimplementowana, stanowią wbudowaną w metodyki zwinne strategię umniejszania ryzyka [1]. Z kolei w [4], choć zarządzanie ryzykiem jest wspominane, to jedyne konkretne rady tam zawarte sprowadzają się do wykonywania częstych audytów oraz poprawienia możliwości śledzenia zmian w kodzie źródłowym w przypadku, gdy błędy w produkcie mogą prowadzić do katastrofalnych skutków w trakcie użytkowania. W [5] wprost stwierdzono, że „analiza ryzyka prowadzona jako osobny proces wydaje się nadmiarem”. Najczęściej jako uzasadnienie zbędności zarządzania ryzykiem podaje się istnienie takich praktyk, jak stały kontakt z klientem, krótkie iteracje czy codzienne spotkania członków zespołu.

3.1. Stały kontakt z klientem

Dzięki stałemu kontaktowi z klientem, wszelkie nieporozumienia, problemy z wymaganiami czy trudności technologiczne zostaną wykryte bardzo szybko. Pozwala to rozwiązać wątpliwości co do cech produktu. Częste spotkania z klientem pozwalają również zmniejszyć ryzyko niepoprawnych estymacji pracochłonności [9]. Jednak istnieje wiele zagrożeń, których źródłem może być klient, a które celowo nie będą przez niego ujawniane – przykładem może być ryzyko utraty płynności finansowej przez klienta. Poza tym wraz z oddawaniem kolejnych gotowych części systemu, pojawia się coraz więcej rodzajów ryzyka, mających źródło w sprzeczności nowo uzyskanych wymagań z tymi już zrealizowanymi lub z niemożnością zrealizowania tychże wymagań w obecnej architekturze.

Metodyki zwinne zakładają, że źródłem wymagań jest użytkownik, ale co jeśli nie zna on w pełni wszystkich uwarunkowań, np. prawnych? Wówczas na liście wymagań może pojawić się pozycja, której implementacja generowałaby tylko i wyłącznie straty – wymagania z listy są odbierane przez zespół projektowy i nikt nie może zagwarantować, że jego członkowie znają dostatecznie prawo ani że dostrzegą konieczność skorzystania z rady eksperta. Po zaimplementowaniu takiego wymagania, funkcjonalność znów odbiera tenże klient, który był źródłem niepoprawnego wymagania. Takie nieprawidłowości mogą być wykryte w produkcie na długo po oddaniu zainfekowanej funkcjonalności, czyli pojawia się ryzyko spowodowane brakiem weryfikacji poprawności wymagań przez osoby inne niż klient i członkowie zespołu projektowego.

Działania wymuszone przez wszelkie zagrożenia są podejmowane dopiero wtedy, gdy pojawiają się pierwsze symptomy spełniania się zagrożenia. Dopiero wtedy będą opracowywane kompleksowe rozwiązania, pozwalające wybrnąć z trudnej sytuacji. Może się to wiązać z opóźnieniami w realizacji projektu lub ze zbędnymi kosztami ponoszonymi przez przedsiębiorstwo na projekt w sytuacji, gdy jego kontynuowanie nie będzie dawało szans na sukces. Wcześniejsze zastanowienie się nad ryzykiem w projekcie umożliwiłoby szybszą diagnozę jego urzeczywistnienia i skróciłyby czas reakcji na nie.

3.2. Częste iteracje

Częste iteracje umożliwiają skorygowanie zakresu projektu czy harmonogramu. Wszelkie nieprawidłowości uwidaczniają się o wiele szybciej i przez to ich negatywny wpływ na projekt jest niewielki. Jednak co, jeśli taka nieprawidłowość nie będzie jednorazowa i krótkotrwała, lecz będzie powoli pogłębiać się przez kolejne iteracje?

Przykładem może tutaj być fluktuacja kadry. Naturalnym momentem do zastanowienia się nad zagrożeniami i podejmowania działań korygujących w wymienionych obszarach jest podsumowanie i planowanie iteracji. Jeśli w pojedynczej iteracji odejdzie z projektu jeden doświadczony, cenny pracownik to, choć będzie to stratą dla zespołu projektowego, nie będzie to jednak sygnałem do podejmowania nagłych, drastycznych kroków naprawczych – obecne w metodykach zwinnych mechanizmy współpracy i dzielenia wiedzy (na przykład programowanie w parach) pozwalają zmniejszyć ryzyko sytuacji, w której tylko jedna osoba posiada unikalną i niezwykle ważną wiedzę. Jednak jeśli tendencja się utrzyma i w kolejnych iteracjach odchodzić będą kolejni wartościowi pracownicy, to sytuacja zyska na powadze, a projekt będzie zagrożony. Częste iteracje w takiej sytuacji nie pomagają, a wręcz utrudniają rozpoznanie długotrwałego ryzyka, gdyż członkowie rozpatrując zagrożenia ograniczają się do jednej iteracji.

Znów rozpatrzenie ryzyka wystąpienia wysokiej fluktuacji kadry zanim jeszcze rozpocznie się projekt (np. poprzez analizę poczynąń najbliższej konkurencji na rynku pracy, ich planów, ofert składanych pracownikom czy też przeprowadzenie badań sprawdzających poziom zadowolenia pracowników) pozwoliłoby na podjęcie odpowiednich kroków zanim jeszcze ryzyko wystąpi.

3.3. Codzienne spotkania członków zespołu

Codzienne spotkania członków zespołu są bardzo wartościowe z punktu widzenia motywacji, budowania zespołu oraz komunikacji wewnątrz niego. Pozwalają nawiązać bezpośredni kontakt, zdać zespołowi sprawozdanie jakie prace zostały wykonane od ostatniego spotkania, a jakie są planowane na bieżący dzień. Ogłaszając swoje plany, członkowie zespołu zobowiązują się do dostarczenia wyników swojej pracy przed kolejnym spotkaniem, co pozwala to innym członkom zespołu efektywniej rozplanować swój czas.

Spotkania te służą także poinformowaniu innych członków zespołu o swoich problemach w realizacji zadań i zwłaszcza ten element może być postrzegany jako pomocny w walce z ryzykiem, lecz i tutaj ryzyko ujawnia się dopiero wtedy, gdy staje się rzeczywistością. W przypadku zagrożenia technologicznego, np. nieświadomego lub nieujawnionego braku odpowiednich kompetencji pracownika odpowiedzialnego za realizację kluczowego zadania, mogą w najlepszym przypadku wystąpić niewielkie opóźnienia w projekcie, a w najgorszym konieczność zmiany technologii. Wpływ na projekt z pewnością nie byłby pomijalny.

3.4. Progresywna redukcja ryzyka

Progresywna redukcja ryzyka jest dość popularnym rozwiązaniem, zaakceptowanym przez wielu autorów publikacji z dziedziny metodyk zwinnych – przykładem metodyki, w której ta praktyka występuje może być SCRUM [8] czy Agile Project Management [5]. Metoda ta polega na rozważaniu podczas planowania nie tylko wartości biznesowej, ale również ryzyka projektu. Planowanie w metodykach zwinnych sprowadza się w znacznym stopniu do wyboru spisanych wcześniej przez klienta wymagań, które mają być zaimplementowane w następnej iteracji. Progresywna redukcja ryzyka oznacza wybór w pierwszej kolejności tych wymagań, które są obciążone większym ryzykiem. Dzięki temu zminimalizowane zostaną koszty spowodowane wystąpieniem ryzyka. Próby zmniejszenia prawdopodobieństwa ryzyka nie są tutaj zawarte.

Czy jednak podejście to pozwala na zidentyfikowanie wszystkich niebezpieczeństw w projekcie? Nie wszystkie rodzaje ryzyka są bezpośrednio związane z wytwarzaniem funkcjonalności, np. ryzyko prawne czy organizacyjne – jeśli proces zarządzania ryzykiem nie wymusza zastanowienia się nad takimi zagrożeniami, bardzo wiele z nich zostanie przeoczonych. Ryzyko może być powiązane nie tylko z poszczególnymi wymaganiami klienta, lecz również z relacjami pomiędzy wymaganiami. W takiej sytuacji ryzyko to nie zostanie wykryte najszybciej jak to możliwe, np. podczas planowania implementacji pierwszego z wymagań obciążonych tymże ryzykiem, ale niespodziewanie pojawi się problem z chwilą rozpoczęcia pracy nad kolejnym powiązaniem wymaganiem, być może dopiero w późniejszych iteracjach.

Również podczas podsumowania kluczowej jednostki procesu wytwarzania oprogramowania, jaką jest pojedyncza iteracja, ryzyko nie jest bezpośrednio adresowane ani kompleksowo analizowane.

Można zatem stwierdzić, że zastosowanie metodyki zwinnej sprawia, że wytwarzanie oprogramowania jest sterowane ryzykiem, jednak nie pozwala w pełni zidentyfikować wszystkich zagrożeń ani tym bardziej nimi zarządzać.

3.5. Źródła ryzyka a metodyki zwinne

Najczęstszymi źródłami niepewności i ryzyka w projektach informatycznych są: wymagania, współpraca z użytkownikami i innymi systemami, zmiany w środowisku projektu, konflikty, jakość zarządzania projektem, zasoby, łańcuch dostaw, polityka, innowacyjność oraz skala projektu [1]. Pierwsze trzy źródła z tej nieuporządkowanej listy są obrane za cel przez metodyki zwinne i ryzyko z nich wynikające jest w znacznym stopniu umniejszone. Ewaluacja

procesu wytwarzania oprogramowania po każdej iteracji sprawia, iż również jakość zarządzania projektem staje się mniej groźnym źródłem ryzyk. Zagrożenia z pozostałych źródeł nie są jednak niwelowane przez metodyki zwinne.

Pewne rodzaje ryzyka są specyficzne dla zwinnego zarządzania projektem lub w podejściu tym znacznie bardziej prawdopodobne. Część z nich powiązana jest z trudnością w skalowaniu metodyk zwinnych przy jednoczesnej chęci zachowania możliwości kontrolowania poprawności kierunku rozwoju produktu poprzez dostarczanie rozszerzanych wersji produktu w stałych, krótkich interwałach czasowych [7]. Ciągła integracja i częste wdrożenia są konieczne, aby odnieść sukces projektu z metodyką zwinną. Jednak ich uzyskanie jest coraz trudniejsze w miarę zwiększania złożoności systemu i rozrostu zespołu projektowego. Nie należy bagatelizować tego problemu, zwłaszcza że skala projektu jest jednym z głównych źródeł zagrożeń.

Kolejna grupa ryzyka wynika z dużej swobody danej programistom. W naturalny sposób preferują oni rozwijanie tych funkcjonalności, które są dla nich „atrakcyjne”, nawet jeśli nie są potrzebne i jeśli nie są uwzględnione w obowiązującej architekturze oprogramowania [7]. Wiele rodzajów ryzyka może być pokłosiem braku pewności, że utworzone fragmenty oprogramowania są stabilne i nie będą już ulegały zmianom.

Zagrożenia są związane również z ograniczeniami obecnymi podczas tworzenia architektury oprogramowania – stałe zmiany wymagań, które pociągają za sobą konieczność częstych modyfikacji architektury, mogą w miarę upływu czasu spowodować jej niespójności oraz wymusić modyfikacje już gotowych elementów. Poziom szczegółowości architektury musi być stosunkowo niski, aby jej dostosowywanie było możliwe i nie pochłaniało zbyt wiele wysiłku. Daje to programistom szerokie pole manewru do interpretacji.

Ponadto, trudne może okazać się przekonanie programistów do bycia wiernym opracowanej architekturze, zwłaszcza gdy ulega ona stałym modyfikacjom. Odstępstwo od architektury może rodzić problemy, nawet jeśli dopuści się go wybitny specjalista techniczny – decyzje przez niego podjęte mogą nie być trafne, ponieważ nie ma on całościowego poglądu na produkt, lecz patrzy jedynie na aktualnie budowany fragment.

W metodykach zwinnych nie ma silnego menedżera projektu, więc istnieje ryzyko znalezienia się w paraliżu decyzyjnym. Niektóre decyzje muszą być podejmowane przez cały zespół, a trudno oczekiwać by każdy zespół przez całe swoje życie był w etapie synergii. W razie konfliktu mogą powstać prze-stoje lub całkowity paraliż pracy, zaś decyzje podejmowane przez osoby o zbyt małym autorytecie formalnym czy finansowym mogą być ignorowane przez członków zespołu.

4. Proponowana metoda zarządzania ryzykiem

Zaproponowana tu metoda jest skonstruowana w taki sposób, aby proces zarządzania ryzykiem był zintegrowany z procesami zarządzania projektem i cyklem życia projektu, wprowadzając minimalny nakład dodatkowej pracy dla członków zespołu. Poczyniono również starania, by modyfikacje procesu wytwarzania oprogramowania nie pozbawiły projektu zwinności i lekkości. Do przedstawienia opracowanej metody zarządzania ryzykiem posłużył model zwinnego wytwarzania oprogramowania (sekcja 2.3).

Za punkt odniesienia przyjęte są tutaj wskazówki zawarte w PMBoK Guide [6], które mówią o następujących procesach zarządzania ryzykiem: planowanie zarządzania ryzykiem, identyfikacja ryzyka, jakościowa analiza ryzyka, ilościowa analiza ryzyka, planowanie reakcji na ryzyko, monitorowanie i kontrola ryzyka. Wszystkie te procesy, poza ilościową analizą ryzyka, znajdują swoje odzwierciedlenie w zaprezentowanych niżej praktykach, składających się na proponowaną metodę zarządzania ryzykiem. Ilościowa analiza ryzyka została pominięta ze względu na jej pracochłonność. Wprowadzenie tak kosztownego procesu kolidowałoby z podstawowymi założeniami podejścia zwinnego do zarządzania projektami.

4.1. Artefakty

Proponowana metoda wprowadza dokument, przeznaczony do przechowywania informacji o ryzyku – rejestr ryzyka. Rejestr ten zawiera opisy wszystkich zidentyfikowanych rodzajów ryzyka obejmujące: identyfikator/nazwę i opis ryzyka, kategorię ryzyka, przyczyny ryzyka, symptomy, planowane reakcje na ryzyko i priorytet (wagę) ryzyka. Dokonanie wpisu do rejestru ryzyka jest poprzedzone identyfikacją ryzyka, analizą jakościową i planowaniem reakcji na nie. Planowanie reakcji może być odroczone, zwłaszcza w przypadku, gdy ryzyko nie jest wysokie lub jest odległe w czasie.

Identyfikator ryzyka ma na celu ułatwienie komunikacji podczas zarządzania tym ryzykiem. Kategoria ryzyka ma informować, jakiego rodzaju ryzyko jest opisywane. Jako przyczyny ryzyka wymienione są zdarzenia i uwarunkowania, które w efekcie mogą skutkować urealnieniem się danego ryzyka. Symptomy to lista objawów, która będzie pomocna we wczesnej detekcji ryzyka. Planowane reakcje na ryzyko zawierają wiele możliwych działań, które mogą być podjęte, gdy ryzyko zacznie się urzeczywistniać. Priorytet ryzyka jest wynikiem jakościowej analizy ryzyka i uwzględnia dotkliwość ryzyka, jego prawdopodobieństwo oraz zdolność zespołu do jego wczesnego wykrycia.

Proponowana metoda wprowadza również zmiany do istniejących artefaktów w metodykach zwinnych, a dokładnie do listy wymagań – w XP będzie to zbiór opowiastek użytkownika, w SCRUM’ie będzie to rejestr wymagań oraz rejestr cyklu. Mianowicie wymagania zostają powiązane ze zidentyfikowanymi rodzajami ryzyka – jedno zidentyfikowane ryzyko może być przypisane do jednego, wielu lub żadnego wymagania; dalsze objaśnienia znajdują się w sekcji 5.5.

4.2. Role

Wprowadzona jest jedna nowa rola do zespołu projektowego – opiekun procesu zarządzania ryzykiem. Zadaniem osoby ją pełniącej jest śledzenie poprawności i staranności zarządzania ryzykiem oraz, w razie konieczności, stymulowanie zespołu do poprawy. Rola ta nie jest na sztywno powiązana z jakąkolwiek formalną funkcją w zespole (na przykład rolą przywódcy zespołu) ani też z żadną inną rolą, zdefiniowaną przez konkretną metodykę zwinną. Zgodnie z ideą samoorganizujących się zespołów, to zespół wskazuje osobę pełniącą tę rolę, a także decyduje, czy jest to rola stała czy rotacyjna oraz czy może ona być współdzielona przez więcej niż jednego członka zespołu. Podział tej roli na więcej niż jedną osobę zalecany jest jedynie w przypadku zbyt dużej ilości ryzyka, które musi być śledzone z należytą uwagą, przy czym opiekun procesu zarządzania ryzykiem jest odpowiedzialny za nadzorowanie jedynie tego ryzyka, które nie są bezpośrednio powiązane z żadnym wymaganiem.

Monitorowanie ryzyka będącego w bezpośredniej relacji z wymaganiami leży w gestii członków zespołu, którzy wykonują zadania wynikające z tych wymagań. W razie zaobserwowania symptomów lub stwierdzenia potrzeby zmiany parametrów ryzyka osoba odpowiedzialna za jego śledzenie informuje o tym zespół na codziennym spotkaniu.

Za zarządzanie ryzykiem odpowiedzialny jest cały zespół. W ramach obowiązujących w podejściu zwinnym samoorganizujących się zespołów, możliwe jest oddelegowanie części lub wszystkich zadań związanych z zarządzaniem ryzykiem do podgrupy zespołu. Jednak odpowiedzialność zawsze spoczywa na całym zespole, co jest w pełni zgodne z zasadami obowiązującymi w metodykach zwinnych. Zatem w optymalnej sytuacji trud identyfikacji, klasyfikacji, analizy, planowania reakcji i podejmowania odpowiednich kroków w odpowiedzi na ryzyko podejmowany byłby przez cały zespół.

4.3. Kwantyfikacja ryzyka

W proponowanej metodzie określone są trzy parametry liczbowe określające dotkliwość konsekwencji, prawdopodobieństwo i zdolność do wczesnego wykrycia ryzyka. Parametry te powinny być wyznaczone przez cały zespół, gdyż to on odpowiada za całościowe zarządzanie ryzykiem. Doskonale do tego celu nadaje się często używana w zarządzaniu zwinnym metoda estymacji – Planning Poker [11].

Dla każdego parametru ryzyka przewidziana jest sztywna skala wartości. Szczegóły przyjętej skali zależą od decyzji zespołu, powinna być tak dobrana, aby wraz ze wzrostem prawdopodobieństwa lub dotkliwości ryzyka, rosła wartość ryzyka oraz by wartość ta malała wraz ze wzrostem zdolności do wczesnego wykrycia ryzyka. Ryzyko może być wykluczone z monitorowania, jeśli całkowita wartość ryzyka nie przekroczy uzgodnionej przez zespół wartości progowej lub jeśli zdolność do wczesnego wykrycia jest niemal zerowa (sekcja 5.5).

4.4. Inicjowanie procesu zarządzania ryzykiem

Planowanie zarządzania ryzykiem powinno być jednym z pierwszych działań w projekcie, przeprowadzone jeszcze zanim zostanie zbudowany wizja produktu. W proponowanej metodzie obejmuje ono zdefiniowanie kategorii ryzyka, które później będą wykorzystywane przy identyfikacji zagrożeń oraz zdefiniowanie skal dopuszczalnych wartości parametrów używanych do kwantyfikacji ryzyka. Dodatkowo ustalone i wyjaśnione powinny być w tym momencie wszelkie inne kwestie, które są konieczne, aby móc rozpocząć zarządzanie ryzykiem według proponowanej metody. Wskazane jest, aby przed wykonaniem planowania zarządzania ryzykiem wyznaczona była wstępnie osoba do pełnienia roli opiekuna procesu zarządzania ryzykiem.

Po znalezieniu początkowych wymagań do produktu, a przed planowaniem edycji produktu wydawanych w kolejnych iteracjach, przeprowadzane jest wyszukiwanie ryzyka. Zalecaną metodą jest burza mózgów z udziałem przedstawiciela klienta oraz wszystkich członków zespołu (lub jego reprezentatywną podgrupą w przypadku zbyt licznej grupy). W razie braku w tym gronie kompetencji z obszarów będących najpoważniejszymi źródłami zagrożeń na spotkanie powinni być zaproszeni również odpowiedni eksperci. Rezultatem jest początkowa lista ryzyka zawarta w rejestrze ryzyka. Po utworzeniu identyfikatora i opisu ryzyka dokonywana jest estymacja ryzyka.

4.5. Planowanie iteracji

Podczas planowania iteracji, zespół powinien prześledzić wszystkie rodzaje ryzyka znajdujące się w rejestrze ryzyka i dla każdego z nich zbadać, czy wystąpiły symptomy, czy ryzyko nadal jest aktualne, czy należy zmodyfikować wpis w rejestrze (np. dokonać ponownej estymacji) oraz czy nie pojawiły się nowe rodzaje ryzyka. Jeżeli zdolność do wczesnego wykrycia ryzyka będzie oceniona jako nikła, wówczas aby uniknąć zbędnych narzutów pracy, zespół może podjąć decyzję o traktowaniu ryzyka jak niepewności – ryzyko takie nie byłoby śledzone w trakcie trwania iteracji. Działanie takie byłoby w tradycyjnych metodykach niezwykle niebezpieczne, jednak metodyki zwinne na tyle wcześnie wykrywają niepewności po ich zmaterializowaniu, że jest to dopuszczalne, o ile dotkliwość i prawdopodobieństwo ryzyka nie sugerują katastrofalnych skutków zaniedbania tego zagrożenia.

Ponadto, przy wyborze wymagań do realizacji w bieżącej iteracji odpowiednia jest ich analiza pod kątem zidentyfikowanych wymagań. Jeśli ryzyko może bezpośrednio zakłócić realizację danego wymagania, to adekwatna informacja wpisywana jest przy tym wymaganiu do rejestru iteracji. Sumaryczna wartość ryzyka powiązanego z danym wymaganiem może być interpretowana jako ryzyko tego wymagania – wartość ta może być wykorzystana do progresywnej redukcji ryzyka.

Analogiczna analiza wykonywana jest dla zadań wynikających z każdego wymagania – jeśli realizacja któregośkolwiek zadania może być dotknięta przez ryzyko powiązane ze źródłowym wymaganiem, to pojawia się odpowiednia adnotacja w rejestrze cyklu. Również tutaj sumaryczna wartość ryzyka dotyczącego zadanie stawałaby się ryzykiem zadania, które byłoby pomocne w ustalaniu kolejności realizacji zadań – zadania obarczone większym ryzykiem powinny być wykonywane jako pierwsze. Osoba (lub para) odpowiedzialna za realizację danego zadania przyjmuje na siebie również odpowiedzialność za monitorowanie ryzyka powiązanego z tym zadaniem.

Suma wszystkich rodzajów ryzyka powiązanych z zadaniami powinna zgadzać się z listą ryzyka przy wymaganiu, z którego te zadania wynikają – można zastosować tutaj podejście zstępujące lub wstępujące do identyfikacji powiązań między ryzykami a zadaniami i wymaganiami.

4.6. Codzienne spotkania

Codzienne spotkania nie tracą nic z dotychczasowego kształtu, jednak dodane są nowe treści wynikające z nowych obowiązków członków zespołu. Podczas spotkania zespół jest także informowany, czy wystąpiły symptomy

nadzorowanych rodzajów ryzyka oraz czy konieczna jest ponowna ocena któregoś z nich, wraz z uzasadnieniem. Również w razie zidentyfikowania nowego ryzyka, członek zespołu powinien poinformować o tym zespół na najbliższym spotkaniu.

Po otrzymaniu jednej z wyżej wymienionych informacji dotyczącej ryzyka, zespół jest zobligowany do przeprowadzenia odpowiednich akcji tego samego dnia, czyli do rozpoczęcia działań w reakcji na ryzyko (korzystając ze wskazówek znajdujących się w rejestrze ryzyk) lub oceny/powtórnej oceny ryzyka.

4.7. Spotkanie retrospektywne

Wraz z zakończeniem iteracji dokonywana jest ewaluacja procesu zarządzania ryzykiem. Każdy członek zespołu, w tym opiekun procesu zarządzania ryzykiem, jak również przedstawiciel klienta mogą zgłosić swoje zastrzeżenia, po czym zespół może zatwierdzić zmiany do procesu, wprowadzając je w życie począwszy od następnej iteracji.

Podsumowanie

Dzięki podstawowym praktykom w metodykach zwinnych, można dość skutecznie uporać się z negatywnymi wpływami niepewności na projekt. Dlatego też zwinne zarządzanie projektami kwestionuje niezbędność procesu zarządzania ryzykiem, wskazując na dopasowanie tego podejścia do projektów prowadzonych w warunkach wysokiego ryzyka. Choć zarządzanie ryzykiem nie jest odradzane we wszystkich metodykach zwinnych, to brakuje w nich elementów, które skłaniałyby zespół do wprowadzenia systematycznego procesu. „Czarnowidztwo”, które jest konieczne do skutecznego zarządzania ryzykiem, nie jest ani naturalnym zachowaniem większości ludzi, ani nie jest dobrze widziane w zespole. Dlatego też nie wystarczy nie zakazywać uprawiania zarządzania ryzykiem, aby ono się pojawiło w metodykach zwinnych – konieczne są wyraźne zalecenia, a czasami nawet nakazy.

Większość praktyk istniejących w metodykach zwinnych skupia się na szybkim wykryciu ryzyka, gdy ono już wystąpi, a te które przyczyniają się do mityzacji, nie obejmują wszystkich możliwych rodzajów zagrożeń. Brakuje praktyk, które jawnie starałyby się antycypować przy uwzględnieniu wszystkich możliwych źródeł ryzyka. Mechanizmy wbudowane w te metodyki nie próbują nawet odpowiedzieć na pytanie, jakie przewidywalne zagrożenia występują w projekcie i jakie jest prawdopodobieństwo ich wystąpienia. Oznacza to traktowanie wszystkich zagrożeń jednakowo – jako nieprzewidywalnych i nie-

poznawalnych z wyprzedzeniem. Wszystkie są rozpatrywane z tą samą skrupulatnością – tyle samo wysiłku będzie poświęcone na zwalczanie małoistotnego zagrożenia, co na zagrożenie krytyczne. Wszystkie te spostrzeżenia wskazują, że metodyki zwinne nie są w stanie zastąpić procesu zarządzania ryzykiem.

Pokazano nie tylko, że wprowadzenie procesu zarządzania ryzykiem do projektów zarządzanych zwinnie jest konieczne, ale również, że jest to możliwe. Zaproponowana metoda zarządzania ryzykiem jest w pełni zintegrowana z metodykami zwinnymi. Poczyniono także starania, aby podstawowe zasady metodyk zwinnych były widoczne również w opracowanej metodzie, a nawet zastosowano niektóre rozwiązania popularne w zarządzaniu zwinnym.

Literatura

1. DeMarco T., Lister T. (2003). *Waltzing with Bears: Managing Risk on Software Projects*. Dorset House.
2. Pritchard C.L. (2001). *Zarządzanie ryzykiem w projektach: Teoria i praktyka*. WIG-Press, Warszawa.
3. Schwaber K., Beedle M. (2001). *Agile Software Development with Scrum*. Prentice Hall, New Jersey.
4. Beck K., Anderes C. (2005). *Extreme Programming Explained: Embrace Change*. Addison-Wesley, Boston.
5. Highsmith J. (2007). *APM: Agile Project Management: Jak tworzyć innowacyjne produkty*. PWN, Warszawa.
6. Nyfjord J., Kajko-Mattsson M. (2008). *Outlining a Model Integrating Risk Management and Agile Software Development*. Euromicro Conference Software Engineering and Advanced Applications.
7. Boehm B., Turner R. (2003). *Using Risk to Balance Agile and Plan-Driven Methods*. IEEE Computer Society.
8. Barton B. (2009). *All-Out Organizational Scrum as an Innovation Value Chain*. Hawaii International Conference on System Sciences.
9. Moløkken-Østfold K., Furulund K.M. (2007). *The Relationship between Customer Collaboration and Software Project Overruns*. AGILE', 72-83.
10. Coram M., Bohner S. (2005). *The Impact of Agile Methods on Software Project Management*. IEEE Conference and Workshops on the Engineering of Computer-Based Systems (ECBS'05).
11. Cohn M. (2005). *Agile Estimating and Planning*. Prentice Hall, New Jersey.

TEORIA GIER I NEGOCJACJI

Hanna Bury

Dariusz Wagner

Instytut Badań Systemowych PAN w Warszawie

WPŁYW WYBORU METODY WYZNACZANIA OCENY GRUPOWEJ NA WYNIK EKSPERTYZY

Wprowadzenie

Z literatury (przedmiotu) wiadomo, że przy wyznaczaniu oceny grupowej bądź wyniku głosowania rezultat w istotny sposób zależy od zastosowanej metody. Stwierdzenie to uzasadnia konieczność dokonania analizy zagadnienia na ile otrzymany wynik odzwierciedla opinie ekspertów (wolę wyborców) czy raczej cechy zastosowanej metody agregacji ocen ekspertów.

Jeżeli nie ma narzuconych ograniczeń odnośnie wyboru metody oceny grupowej, w praktyce może okazać się, że rozwiązania uzyskane przy użyciu różnych metod są zgodne lub zbliżone albo całkowicie rozbieżne. W pierwszej sytuacji można przyjąć, że ocena grupowa wiernie odzwierciedla opinie ekspertów (wolę wyborców). W szczególnym przypadku różne metody mogą wskazywać ten sam obiekt jako zwycięzcę. Baharad i Nitzan [1; 2] formułują warunki, jakie powinny być spełnione, aby zapewnić zgodność wyboru najlepszego obiektu wyznaczonego za pomocą różnych metod.

Można również sformułować zadanie wyboru takiej metody wyznaczania oceny grupowej, która zapewni zgodność otrzymanego rozwiązania z założonymi wcześniej warunkami.

Saari [4] przeanalizował zagadnienie wyznaczania dla danego zestawu ocen ekspertów wszystkich możliwych rozwiązań uzyskanych w wyniku zastosowania różnych metod pozycyjnych. W omawianym opracowaniu jest podana szczegółowa analiza przypadku trzech obiektów. W pracy podejście to będzie rozszerzone na przykład czterech obiektów.

1. Podstawowe definicje

1.1. Zadanie wyznaczania oceny grupowej

Zakładamy, że dany jest zbiór n obiektów $\mathcal{O} = \{O_1, O_2, \dots, O_n\}$ oraz zbiór K ekspertów, których zadaniem jest uporządkowanie tego zbioru zgodnie z przyjętym kryterium (zbiorem kryteriów). Zazwyczaj przyjmuje się, że obiekt uważany za najlepszy – w sensie przyjętego kryterium lub zbioru kryteriów – jest umieszczony na pierwszej pozycji, zaś obiekt uważany za najgorszy zajmuje ostatnią pozycję.

1.2. Metody pozycyjne wyznaczania oceny grupowej

Rozważania zostaną ograniczone do tzw. metod pozycyjnych, tzn. takich, w których ocena obiektu jest wyznaczana na podstawie jego pozycji w uporządkowaniu. Z reguły przyjmuje się – dla uproszczenia rozważań – że w uporządkowaniach podanych przez ekspertów nie występują obiekty równoważne oraz że liczba pozycji we wszystkich uporządkowaniach jest taka sama.

Wektor wag (w teorii głosowania nazywany wektorem głosowania) ma następującą postać

$$\mathbf{w} = (w_1, \dots, w_n) \in R_n, \quad (1)$$

przy czym waga w_j – jest to liczba przyporządkowana obiektowi zajmującemu pozycję j w uporządkowaniu.

Zazwyczaj przyjmuje się, że $\forall_j w_j \geq w_{j+1}$ oraz, że $w_1 > w_n$, $j = 1, \dots, n-1$.

Łączna liczba punktów s_i , jakie uzyskuje dany obiekt O_i dana jest zależnością

$$s_i = \sum_{k=1}^K \sum_{j=1}^n \delta_j^{ik} w_j, \quad \delta_j^{ik} = \begin{cases} 1 & \text{jeżeli obiekt } O_i \text{ zajmuje pozycję } j \\ & \text{w uporządkowaniu podanym przez eksperta } k \\ 0 & \text{w przeciwnym razie} \end{cases} \quad (2)$$

Zwycięzcą zostaje obiekt, dla którego wartość wskaźnika s_i jest największa.

Przyjęcie konkretnej postaci wektora wag prowadzi do określonej metody wyznaczania oceny grupowej, np.:

$w=(1, 0, \dots, 0)$ odpowiada metodzie większości

$w = (\underbrace{1, 1, \dots, 1}_m, 0, \dots, 0)$ odpowiada głosowaniu na m kandydatów

$w=(1, 1, \dots, 1, 0)$ odpowiada metodzie antywiększości

$w=(n-1, \dots, n-j, \dots, 1, 0)$ odpowiada metodzie Bordy.

Ogólnie wektor głosowania można przedstawić jako kombinację wypukłą wektorów głosowania typu „głosuj na m kandydatów”

$(1, 0, \dots, 0), (1, 1, 0, \dots, 0), \dots, (1, 1, 1, \dots, 1, 0), m=1, \dots, n-1.$

Aby zilustrować prowadzone rozważania, rozpatrzmy przykład podany w pracy [6].

Przykład 1

Przyjmijmy, że jedenastu ekspertów podało następujące uporządkowania zbioru trzech obiektów $\{O_1, O_2, O_3\}$:

liczba ekspertów	3	2	2	4
uporząd- kowanie	$O_1 \succ O_2 \succ O_3$	$O_1 \succ O_3 \succ O_2$	$O_2 \succ O_3 \succ O_1$	$O_3 \succ O_2 \succ O_1$

Zastosowanie trzech wybranych metod pozycyjnych prowadzi do następujących wyników:

Metoda/wektor wag	Liczba punktów			Uporządkowanie	Zwycięzca
	O_1	O_2	O_3		
Większości (1, 0, 0)	5	2	4	$O_1 \succ O_3 \succ O_2$	O_1
Antywiększości (1, 1, 0)	5	9	8	$O_2 \succ O_3 \succ O_1$	O_2
Bordy (2, 1, 0)	10	11	12	$O_3 \succ O_2 \succ O_1$	O_3

Przykład ten pokazuje, że nawet w tak prostym przypadku każdy obiekt może zostać zwycięzcą pod warunkiem zastosowania odpowiedniej metody głosowania.

W swoich pracach (między innymi [4; 5; 6; 7]) Saari przeanalizował to zagadnienie szczegółowo, dochodząc do interesujących wniosków, które sformułował w postaci twierdzeń. Niektóre z nich zostaną przytoczone w następnych rozdziałach.

1.3. Definicja profilu

Saari zaproponował, aby oceny podawane przez ekspertów opisywać za pomocą tzw. profilu. Profil jest to pełna lista uporządkowań obiektów podanych przez ekspertów. Określenie profilu wymaga uwzględnienia wszystkich uporządkowań rozpatrywanych obiektów. W rozważanym przypadku, tzn. przy założeniu, że w uporządkowaniach nie występują obiekty równoważne, liczba wszystkich możliwych uporządkowań n obiektów wynosi $n!$. Liczby uporządkowań n obiektów z równoważnościami podano w pracy [3].

Założmy, że p_ℓ oznacza liczbę ekspertów, których opinia ma postać uporządkowania P^ℓ , ($\ell = 1, \dots, n!$). Jeżeli liczba ekspertów, których opinie bierzemy pod uwagę jest równa K , to

$$\sum_{\ell=1}^{n!} p_\ell = K \quad (3)$$

Profilem $\mathbf{p}$ będziemy nazywać wektor o $n!$ składowych p_ℓ .

2. Podejście geometryczne Saariego

W metodach pozycyjnych stosuje się tzw. znormalizowane wektory wag, tzn. takie, których suma składowych wynosi 1. Nie wszystkie wektory wag spełniają ten warunek. Na przykład dla czterech obiektów wektor wag odpowiadający metodzie Bordy ma postać

$$\left(\frac{3}{4}, \frac{2}{4}, \frac{1}{4}, 0 \right) \quad (4)$$

Saari [5] wykazał, że obowiązuje następujące twierdzenie.

Twierdzenie 1

Jeżeli dwa wektory głosowania są powiązane zależnością liniową (a, b skalarne, $a, b > 0$)

$$\mathbf{w}^2 = a\mathbf{w}^1 + b(1, 1, \dots, 1, 0) \quad (5)$$

wówczas oceny grupowe (uporządkowania) uzyskane przy użyciu tych wektorów są takie same.

Mamy zatem

$$w_i^2 = aw_i^1 + b, i=1, \dots, n-1; w_n^2 = 0; w_n^1 = 0 \text{ oraz } \sum_{i=1}^{n-1} w_i^1 = 1 \quad (6)$$

Można wykazać, że

$$w_{n-1}^1 = \frac{W_{n-1} - b}{W - b(n-1)}, \quad w^1 = \frac{(w^2 - b)}{W - b(n-1)} \quad (7)$$

oraz

$$\frac{W_i}{W - b(n-1)} + w_{i+1}^1 = w_i^1 \quad \text{gdzie} \quad W_i = w_i^2 - w_{i+1}^2, \quad W = \sum_{i=1}^{n-1} w_i^2 \quad (8)$$

Wykorzystując wzór (8) możemy – przyjmując określoną wartość współczynnika b – wyznaczyć kolejne wartości w_i^1 .

Przykład 2

Załóżmy, że $n=4$ oraz że wektor wag ma postać $w^2 = \left(\frac{3}{4}, \frac{2}{4}, \frac{1}{4}, 0\right)$

(metoda Bordy). Wykorzystując wzór (6) oraz przyjmując $b = \frac{1}{8}$ otrzymujemy

$$w^1 = \left(\frac{5}{9}, \frac{3}{9}, \frac{1}{9}, 0\right) \quad (9)$$

W przykładzie 3 pokazano, że ocena grupowa dla uporządkowań z przykładu 1 wyznaczona za pomocą metody Bordy przy użyciu wektorów w^2 oraz w^1 jest taka sama.

Twierdzenie 1 upoważnia do przyjęcia znormalizowanej postaci wektora wag (gdzie $w_n = 0$ oraz $\sum_{i=1}^n w_i = 1$) oraz do wprowadzenia przestrzeni znormalizowanych wektorów wag będącej w przypadku n obiektów częścią $n-2$ wymiarowego sympleksu danego zależnością [4]

$$\mathcal{W}^n = \left\{ w^n \mid w_i \geq w_{i+1}, i = 1, \dots, n-1, w_n = 0, \sum_{i=1}^n w_i = 1 \right\} \quad (10)$$

Istotną rolę w podejściu Saariego odgrywają wektory głosowania na s kandydatów E_s^n przedstawione w postaci znormalizowanej

$$E_s^n = \left(\underbrace{\frac{1}{s}, \dots, \frac{1}{s}}_{s \text{ razy}}, 0, \dots, 0 \right), \quad s = 1, \dots, n-1 \quad (11)$$

Są one wierzchołkami sympleksu $\mathcal{W}^n$. Wektory $\{\mathbf{E}_s^n\}_{s=1}^{n-1}$ tworzą bazę wypukłą dla wypukłego zbioru $\mathcal{W}^n$, tak więc każdy wektor głosowania $\mathbf{w}^n \in \mathcal{W}^n$ ma jednoznaczną reprezentację wypukłą

$$\mathbf{w}^n = \sum_{s=1}^{n-1} \lambda_s \mathbf{E}_s^n, \quad \lambda_s \geq 0, \quad \sum_{s=1}^{n-1} \lambda_s = 1 \quad (12)$$

Wektory wypukłych wag w sympleksie $(n-2)$ -wymiarowym

$$\Lambda^n = \left\{ \lambda = (\lambda_1, \dots, \lambda_{n-1}) \mid \lambda_s \geq 0, \quad \sum_{s=1}^{n-1} \lambda_s = 1 \right\} \quad (13)$$

są jednoznacznie powiązane z wektorami głosowania ze zbioru $\mathcal{W}^n$.

Twierdzenie 2 [4]

Znormalizowany wynik głosowania $f(\mathbf{p}, \mathbf{w}^n)$ jest wyznaczany na podstawie zależności

$$f(\mathbf{p}, \mathbf{w}^n) = \sum_{s=1}^{n-1} \lambda_s f(\mathbf{p}, \mathbf{E}_s^n) \quad (14)$$

2.1. Przykład postępowania

Założymy, że w ekspertyzie uczestniczy 17 ekspertów, których zadaniem jest uporządkowanie zbioru czterech obiektów. Przyjmujemy, że uporządkowania podane przez ekspertów można podzielić na cztery grupy: P^1 , P^2 , P^3 , P^4 , przy czym:

5 ekspertów podało uporządkowanie	$P^1: O_1 \succ O_3 \succ O_4 \succ O_2$
2 ekspertów podało uporządkowanie	$P^2: O_2 \succ O_1 \succ O_3 \succ O_4$
4 ekspertów podało uporządkowanie	$P^3: O_3 \succ O_1 \succ O_4 \succ O_2$
oraz 6 ekspertów podało uporządkowanie	$P^4: O_1 \succ O_4 \succ O_2 \succ O_3.$

W rozpatrywanym przypadku na $4! = 24$ składowe profilu $\mathbf{p}$ tylko cztery są różne od zera. Zgodnie z metodologią podaną przez Saariego analizując profil $\mathbf{p}$ utworzony przez uporządkowania $\{P^1, P^2, P^3, P^4\}$ czterech obiektów należy rozpatrzyć $(4-1) = 3$ wektory $\mathbf{E}_s^4$

$$\mathbf{E}_1^4 = (1, 0, 0, 0), \quad \mathbf{E}_2^4 = \left(\frac{1}{2}, \frac{1}{2}, 0, 0\right), \quad \mathbf{E}_3^4 = \left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3}, 0\right) \quad (15)$$

Dla uproszczenia stosujemy zapis $\mathbf{p} = \left(\frac{p^1}{17}, \frac{p^2}{17}, \frac{p^3}{17}, \frac{p^4}{17} \right)$ uwzględniający

tylko niezerowe składowe profilu $\mathbf{p}$.

Następnie należy utworzyć trzy funkcje $f(\mathbf{p}, \mathbf{E}_s^4)$, $s=1, 2, 3$ określające wynik głosowania otrzymany przy użyciu danego wektora wag (czyli konkretnej metody pozycyjnej) do uporządkowań $\{P^1, P^2, P^3, P^4\}$.

Zasada tworzenia wyrażeń $f(\mathbf{p}, \mathbf{E}_s^4)$ jest następująca: dla każdego uporządkowania P^ℓ , $\ell=1, \dots, 4$, składowe wektora $\mathbf{E}_s^4$ zależą od kolejności obiektów w uporządkowaniu. Dla $\mathbf{E}_s^4$, $s=1, 2, 3$, ułamek $\frac{1}{s}$ występuje na pozycjach odpowiadających obiektom zajmującym pierwszych s miejsc w uporządkowaniu.

Dla $\mathbf{E}_1^4$ jedynka występuje na pozycji odpowiadającej obiektowi zajmującemu pierwsze miejsce w uporządkowaniu.

Mamy więc dla $P^1 - \mathbf{E}_1^4 = (1, 0, 0, 0)$, dla $P^2 - \mathbf{E}_1^4 = (0, 1, 0, 0)$,

dla $P^3 - \mathbf{E}_1^4 = (0, 0, 1, 0)$ oraz dla $P^4 - \mathbf{E}_1^4 = (1, 0, 0, 0)$.

Zatem dla $\mathbf{E}_1^4$ mamy

$$\begin{aligned} f(\mathbf{p}, \mathbf{E}_1^4) = & \frac{5}{17} \underbrace{\begin{pmatrix} O_1 & O_2 & O_3 & O_4 \\ 1 & 0 & 0 & 0 \end{pmatrix}}_{\mathbf{p}^1} + \frac{2}{17} \underbrace{\begin{pmatrix} O_1 & O_2 & O_3 & O_4 \\ 0 & 1 & 0 & 0 \end{pmatrix}}_{\mathbf{p}^2} \\ & + \frac{4}{17} \underbrace{\begin{pmatrix} O_1 & O_2 & O_3 & O_4 \\ 0 & 0 & 1 & 0 \end{pmatrix}}_{\mathbf{p}^3} + \frac{6}{17} \underbrace{\begin{pmatrix} O_1 & O_2 & O_3 & O_4 \\ 1 & 0 & 0 & 0 \end{pmatrix}}_{\mathbf{p}^4} = \begin{pmatrix} \frac{11}{17} & \frac{2}{17} & \frac{4}{17} & 0 \end{pmatrix} \end{aligned} \quad (16)$$

czyli uporządkowanie obiektów otrzymane przy użyciu wektora wag $\mathbf{E}_1^4$ ma postać $O_1 \succ O_3 \succ O_2 \succ O_4$.

(17)

Dla $\mathbf{E}_2^4$ ułamek $\frac{1}{2}$ występuje na pozycjach odpowiadających obiektom zajmującym pierwsze dwa miejsca w uporządkowaniu.

Dla P^1 mamy $\mathbf{E}_2^4 = \left(\frac{1}{2}, 0, \frac{1}{2}, 0 \right)$, dla $P^2 - \mathbf{E}_2^4 = \left(\frac{1}{2}, \frac{1}{2}, 0, 0 \right)$, dla $P^3 -$

$\mathbf{E}_2^4 = \left(\frac{1}{2}, 0, \frac{1}{2}, 0 \right)$ oraz dla $P^4 - \mathbf{E}_2^4 = \left(\frac{1}{2}, 0, 0, \frac{1}{2} \right)$. Zatem dla $\mathbf{E}_2^4$ mamy

$$\begin{aligned}
 f(\mathbf{p}, \mathbf{E}_2^4) &= \frac{5}{17} \underbrace{\left(\frac{O_1}{2}, \frac{O_2}{0}, \frac{O_3}{2}, \frac{O_4}{0} \right)}_{\mathbf{p}^1} + \frac{2}{17} \underbrace{\left(\frac{O_1}{2}, \frac{O_2}{2}, \frac{O_3}{0}, \frac{O_4}{0} \right)}_{\mathbf{p}^2} + \frac{4}{17} \underbrace{\left(\frac{O_1}{2}, \frac{O_2}{0}, \frac{O_3}{2}, \frac{O_4}{0} \right)}_{\mathbf{p}^3} + \\
 &\frac{6}{17} \underbrace{\left(\frac{O_1}{2}, \frac{O_2}{0}, \frac{O_3}{0}, \frac{O_4}{2} \right)}_{\mathbf{p}^4} = \left(\frac{5+2+4+6}{34}, \frac{2}{34}, \frac{5+4}{34}, \frac{6}{34} \right) = \left(\frac{O_1}{34}, \frac{O_2}{34}, \frac{O_3}{34}, \frac{O_4}{34} \right)
 \end{aligned} \quad (18)$$

czyli uporządkowanie obiektów otrzymane przy użyciu wektora wag $\mathbf{E}_2^4$ ma postać $O_1 \succ O_3 \succ O_4 \succ O_2$. (19)

Dla $\mathbf{E}_3^4$ ułamek $\frac{1}{3}$ występuje na pozycjach odpowiadających obiektom zajmującym pierwsze trzy miejsca w uporządkowaniu. Mamy więc dla $\mathbf{P}^1$ - $\mathbf{E}_3^4 = \left(\frac{1}{3}, 0, \frac{1}{3}, \frac{1}{3} \right)$, dla $\mathbf{P}^2$ - $\mathbf{E}_3^4 = \left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3}, 0 \right)$, dla $\mathbf{P}^3$ - $\mathbf{E}_3^4 = \left(\frac{1}{3}, 0, \frac{1}{3}, \frac{1}{3} \right)$ oraz dla $\mathbf{P}^4$ - $\mathbf{E}_3^4 = \left(\frac{1}{3}, \frac{1}{3}, 0, \frac{1}{3} \right)$.

Zatem dla $\mathbf{E}_3^4$ mamy

$$\begin{aligned}
 f(\mathbf{p}, \mathbf{E}_3^4) &= \frac{5}{17} \underbrace{\left(\frac{O_1}{3}, \frac{O_2}{0}, \frac{O_3}{3}, \frac{O_4}{3} \right)}_{\mathbf{p}^1} + \frac{2}{17} \underbrace{\left(\frac{O_1}{3}, \frac{O_2}{3}, \frac{O_3}{3}, \frac{O_4}{0} \right)}_{\mathbf{p}^2} + \frac{4}{17} \underbrace{\left(\frac{O_1}{3}, \frac{O_2}{0}, \frac{O_3}{3}, \frac{O_4}{3} \right)}_{\mathbf{p}^3} + \frac{6}{17} \underbrace{\left(\frac{O_1}{3}, \frac{O_2}{3}, \frac{O_3}{0}, \frac{O_4}{3} \right)}_{\mathbf{p}^4} \\
 &= \left(\frac{5+2+4+6}{51}, \frac{2+6}{51}, \frac{5+2+4}{51}, \frac{5+4+6}{51} \right) = \left(\frac{O_1}{51}, \frac{O_2}{51}, \frac{O_3}{51}, \frac{O_4}{51} \right)
 \end{aligned} \quad (20)$$

czyli uporządkowanie obiektów otrzymane przy użyciu wektora wag $\mathbf{E}_3^4$ ma postać $O_1 \succ O_4 \succ O_3 \succ O_2$. (21)

Aby w rozpatrywanym przykładzie wyznaczyć zbiór wszystkich możliwych uporządkowań uzyskanych przy użyciu różnych metod pozycyjnych, należy – zgodnie z metodologią podaną przez Saariego – wyznaczyć funkcję $f(\mathbf{p}, \mathbf{w}_\lambda^4)$, gdzie wektor wag $\mathbf{w}_\lambda^4$ ma postać

$$\mathbf{w}_\lambda^4 = \lambda_1 \mathbf{E}_1^4 + \lambda_2 \mathbf{E}_2^4 + \lambda_3 \mathbf{E}_3^4 \quad (22)$$

przy czym

$$\lambda_1 + \lambda_2 + \lambda_3 = 1 \quad (23)$$

czyli

$$\lambda_3 = 1 - (\lambda_1 + \lambda_2) \quad (24)$$

skąd

$$\mathbf{w}_\lambda^4 = \lambda_1 \mathbf{E}_1^4 + \lambda_2 \mathbf{E}_2^4 + (1 - \lambda_1 - \lambda_2) \mathbf{E}_3^4 \quad (25)$$

Na podstawie (20), (24) mamy

$$\begin{aligned} \mathbf{w}_\lambda^4 &= \lambda_1 (1, 0, 0, 0) + \lambda_2 (1/2, 1/2, 0, 0) + (1 - \lambda_1 - \lambda_2) (1/3, 1/3, 1/3, 0) \\ &= \left(\overbrace{\lambda_1 + \frac{1}{2}\lambda_2 + \frac{1}{3}(1 - \lambda_1 - \lambda_2)}^{j=1}, \overbrace{0 + \frac{1}{2}\lambda_2 + \frac{1}{3}(1 - \lambda_1 - \lambda_2)}^{j=2}, \overbrace{\frac{1}{3}(1 - \lambda_1 - \lambda_2)}^{j=3}, \overbrace{0}^{j=4} \right) \\ &= \left(\frac{\overbrace{6\lambda_1 + 3\lambda_2 + 2 - 2\lambda_1 - 2\lambda_2}^{j=1}}{6}, \frac{\overbrace{3\lambda_2 + 2 - 2\lambda_1 - 2\lambda_2}^{j=2}}{6}, \frac{\overbrace{2 - 2\lambda_1 - 2\lambda_2}^{j=3}}{6}, \overbrace{0}^{j=4} \right) \\ &= \left(\frac{\overbrace{2 + 4\lambda_1 + \lambda_2}^{j=1}}{6}, \frac{\overbrace{2 + \lambda_2 - 2\lambda_1}^{j=2}}{6}, \frac{\overbrace{2 - 2\lambda_1 - 2\lambda_2}^{j=3}}{6}, \overbrace{0}^{j=4} \right) = (w_1^4, w_2^4, w_3^4, w_4^4) \end{aligned} \quad (26)$$

Zapis ten pokazuje, jakie wagi są przyporządkowywane obiektom zajmującym poszczególne pozycje $j=1, 2, 3, 4$. Zatem, mając składowe wektora $\mathbf{w}_\lambda^4$ należy wyznaczyć $f(\mathbf{p}, \mathbf{w}_\lambda^4)$.

Biorąc pod uwagę postać uporządkowań $\{P^1, P^2, P^3, P^4\}$ otrzymamy

$$\begin{aligned} f(\mathbf{p}, \mathbf{w}_\lambda^4) &= \frac{5}{17} \underbrace{(w_1^4, w_4^4, w_2^4, w_3^4)}_{P^1} + \frac{2}{17} \underbrace{(w_2^4, w_1^4, w_3^4, w_4^4)}_{P^2} + \\ &\quad \frac{4}{17} \underbrace{(w_2^4, w_4^4, w_1^4, w_3^4)}_{P^3} + \frac{6}{17} \underbrace{(w_1^4, w_2^4, w_4^4, w_3^4)}_{P^4} \end{aligned} \quad (27)$$

gdzie składowe wektora $\mathbf{w}_\lambda^4$ dane są wzorem (26).

Składowe wektorowego wyrażenia $f(\mathbf{p}, \mathbf{w}_\lambda^4)$ są, jak następuje:

$$\begin{aligned} \text{dla } O_1: & \frac{10 + 5\lambda_2 + 20\lambda_1 + 4 + 2\lambda_2 - 4\lambda_1 + 8 + 4\lambda_2 - 8\lambda_1 + 12 + 6\lambda_2 + 24\lambda_1}{102} \\ & = \frac{34 + 17\lambda_2 + 32\lambda_1}{102} \end{aligned} \quad (28)$$

$$\text{dla } O_2: \frac{4 + 2\lambda_2 + 8\lambda_1 + 12 - 12\lambda_1 - 12\lambda_2}{102} = \frac{16 - 10\lambda_2 - 4\lambda_1}{102} \quad (29)$$

$$\text{dla } O_3: \frac{10 + 5\lambda_2 - 10\lambda_1 + 4 - 4\lambda_2 - 4\lambda_1 + 8 + 4\lambda_2 + 16\lambda_1}{102} = \frac{22 + 5\lambda_2 + 2\lambda_1}{102} \quad (30)$$

$$\begin{aligned} \text{dla } O_4: & \frac{10 - 10\lambda_2 - 10\lambda_1 + 8 - 8\lambda_2 - 8\lambda_1 + 12 + 6\lambda_2 - 12\lambda_1}{102} = \\ & \frac{30 - 30\lambda_1 - 12\lambda_2}{102} \end{aligned} \quad (31)$$

Wyrażenie $f(\mathbf{p}, \mathbf{w}_\lambda^4)$ można też określić w inny sposób. Wiadomo, że [4]

$$\begin{aligned} f(\mathbf{p}, \mathbf{w}_\lambda^4) &= f(\mathbf{p}, \lambda_1 \mathbf{E}_1^4 + \lambda_2 \mathbf{E}_2^4 + \lambda_3 \mathbf{E}_3^4) \\ &= \lambda_1 f(\mathbf{p}, \mathbf{E}_1^4) + \lambda_2 f(\mathbf{p}, \mathbf{E}_2^4) + \lambda_3 f(\mathbf{p}, \mathbf{E}_3^4) \end{aligned} \quad (32)$$

Biorąc pod uwagę, że wyrażenie $f(\mathbf{p}, \mathbf{w}_\lambda^4)$ jest liniowe oraz uwzględniając (20) i (23) mamy

$$\begin{aligned} f(\mathbf{p}, \mathbf{w}_\lambda^4) &= \lambda_1 f(\mathbf{p}, \mathbf{E}_1^4) + \lambda_2 f(\mathbf{p}, \mathbf{E}_2^4) + \lambda_3 f(\mathbf{p}, \mathbf{E}_3^4) \\ &= \lambda_1 f(\mathbf{p}, \mathbf{E}_1^4) + \lambda_2 f(\mathbf{p}, \mathbf{E}_2^4) + (1 - \lambda_1 - \lambda_2) f(\mathbf{p}, \mathbf{E}_3^4) \end{aligned} \quad (33)$$

Uwzględniając postacie wyrażen $f(\mathbf{p}, \mathbf{E}_1^4)$ (16), $f(\mathbf{p}, \mathbf{E}_2^4)$ (18), $f(\mathbf{p}, \mathbf{E}_3^4)$ (20) mamy

$$\begin{aligned} f(\mathbf{p}, \mathbf{w}_\lambda^4) &= \lambda_1 \left(\frac{11}{17}, \frac{2}{17}, \frac{4}{17}, 0 \right) + \lambda_2 \left(\frac{17}{34}, \frac{2}{34}, \frac{9}{34}, \frac{6}{34} \right) \\ &+ (1 - \lambda_1 - \lambda_2) \left(\frac{17}{51}, \frac{8}{51}, \frac{11}{51}, \frac{15}{51} \right) \end{aligned} \quad (34)$$

Poszczególne składowe wektora $f(\mathbf{p}, \mathbf{w}_\lambda^4)$, mają więc postać:

$$\begin{aligned} \text{dla } O_1: \frac{11}{17}\lambda_1 + \frac{17}{34}\lambda_2 + \frac{17}{51}(1-\lambda_1-\lambda_2) &= \frac{66\lambda_1 + 51\lambda_2 + 34 - 34\lambda_1 - 34\lambda_2}{102} \\ &= \frac{34 + 32\lambda_1 + 17\lambda_2}{102} \end{aligned} \quad (35)$$

$$\begin{aligned} \text{dla } O_2: \frac{2}{17}\lambda_1 + \frac{2}{34}\lambda_2 + \frac{8}{51}(1-\lambda_1-\lambda_2) &= \frac{12\lambda_1 + 6\lambda_2 + 16 - 16\lambda_1 - 16\lambda_2}{102} \\ &= \frac{16 - 4\lambda_1 - 10\lambda_2}{102} \end{aligned} \quad (36)$$

$$\begin{aligned} \text{dla } O_3: \frac{4}{17}\lambda_1 + \frac{9}{34}\lambda_2 + \frac{11}{51}(1-\lambda_1-\lambda_2) &= \frac{24\lambda_1 + 27\lambda_2 + 22 - 22\lambda_1 - 22\lambda_2}{102} \\ &= \frac{22 + 2\lambda_1 + 5\lambda_2}{102} \end{aligned} \quad (37)$$

$$\begin{aligned} \text{dla } O_4: 0 \cdot \lambda_1 + \frac{6}{34}\lambda_2 + \frac{15}{51}(1-\lambda_1-\lambda_2) &= \frac{18\lambda_2 + 30 - 30\lambda_1 - 30\lambda_2}{102} \\ &= \frac{30 - 30\lambda_1 - 12\lambda_2}{102} \end{aligned} \quad (38)$$

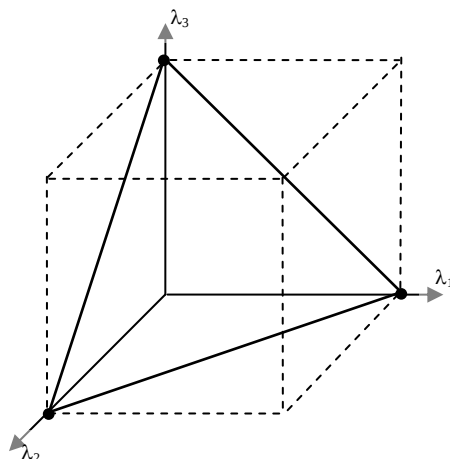
Otrzymany wynik pokrywa się z (28)-(31). Można zatem – w zależności od postaci zadania – stosować jedno z dwu przedstawionych podejść.

Jeżeli współczynniki λ_i spełniają warunek $\sum_{i=1}^3 \lambda_i = 1$, to wartości tych współczynników muszą należeć obszarowi wyznaczonego przez wierzchołki $(1, 0, 0)$, $(0, 1, 0)$, $(0, 0, 1)$ (rys. 1).

$$\text{Jeżeli } \lambda_3 = 0, \text{ to } \lambda_1 + \lambda_2 = 1, \text{ czyli } \lambda_2 = 1 - \lambda_1. \quad (39)$$

$$\text{Jeżeli } \lambda_2 = 0, \text{ to } \lambda_1 + \lambda_3 = 1, \text{ czyli } \lambda_3 = 1 - \lambda_1. \quad (40)$$

$$\text{Jeżeli } \lambda_1 = 0, \text{ to } \lambda_2 + \lambda_3 = 1, \text{ czyli } \lambda_3 = 1 - \lambda_2. \quad (41)$$



Rys. 1. Obszar dopuszczalnych wartości λ_i , $i=1, 2, 3$

W celu ustalenia, jakiego typu uporządkowania można uzyskać jako ocenę grupową dla danego profilu $\mathbf{p}$ za pomocą różnych metod pozycyjnych, należy przebadac wpływ zmienności parametrów λ_1 i λ_2 na wartość składowych wektora $f(\mathbf{p}, \mathbf{w}_\lambda^4)$.

$$1. \text{ Warunek } O_1 \succ O_2 \text{ oznacza, że } 34 + 32\lambda_1 + 17\lambda_2 > 16 - 4\lambda_1 - 10\lambda_2, \quad (42)$$

$$\text{czyli } 18 > -36\lambda_1 - 10\lambda_2. \quad (43)$$

Ponieważ $\lambda_1 \geq 0$ i $\lambda_2 \geq 0$ warunek ten jest zawsze spełniony.

$$2. \text{ Warunek } O_2 \succ O_3 \text{ oznacza, że } 16 - 4\lambda_1 - 10\lambda_2 > 22 + 2\lambda_1 + 5\lambda_2, \quad (44)$$

$$\text{czyli } -6\lambda_1 - 15\lambda_2 > 6. \quad (45)$$

Warunek ten nie może być spełniony, ponieważ $\lambda_1 \geq 0$ i $\lambda_2 \geq 0$.

$$3. \text{ Warunek } O_3 \succ O_4 \text{ oznacza, że } 22 + 2\lambda_1 + 5\lambda_2 > 30 - 30\lambda_1 - 12\lambda_2, \quad (46)$$

$$\text{czyli } 4\lambda_1 + \frac{17}{8}\lambda_2 > 1. \quad (47)$$

$$\text{Dla } \lambda_1 = 0 \text{ mamy } \lambda_2 > \frac{8}{17}, \text{ dla } \lambda_1 = 1 \text{ mamy } 4 + \frac{17}{8}\lambda_2 > 1. \quad (48)$$

Warunek ten jest zawsze spełniony, ponieważ $\lambda_2 \geq 0$.

$$4. \text{ Warunek } O_4 \succ O_3 \text{ oznacza, że } 30 - 30\lambda_1 - 12\lambda_2 > 22 + 2\lambda_1 + 5\lambda_2, \quad (49)$$

$$\text{czyli } 1 > 4\lambda_1 + \frac{17}{8}\lambda_2. \quad (50)$$

$$\text{Dla } \lambda_2 = 0 \text{ mamy } \lambda_1 < \frac{1}{4}, \text{ dla } \lambda_1 = 0 \text{ mamy } \lambda_2 < \frac{8}{17}. \quad (51)$$

5. Warunek $O_3 \approx O_4$.

Z porównania warunków 3 i 4 wynika, że równoważność obiektów O_3 i O_4 otrzymujemy wtedy, gdy spełniona jest równość $\lambda_2 = \frac{8}{17} - \frac{32}{17} \lambda_1$. (52)

6. Warunek $O_2 \succ O_4$ oznacza, że $16 - 4\lambda_1 - 10\lambda_2 > 30 - 30\lambda_1 - 12\lambda_2$ (53)

$26\lambda_1 + 2\lambda_2 > 14$, czyli $\lambda_2 > 7 - 13\lambda_1$. (54)

Jeżeli $\lambda_2 = 0$, to $\lambda_1 > \frac{7}{13}$. (55)

7. Warunek $O_4 \succ O_2$ oznacza, że $30 - 30\lambda_1 - 12\lambda_2 > 16 - 4\lambda_1 - 10\lambda_2$, (56)

czyli $\lambda_2 < 7 - 13\lambda_1$. (57)

Jeżeli $\lambda_2 = 0$, to $\lambda_1 < \frac{7}{13}$. (58)

8. Warunek $O_2 \approx O_4$.

Z porównania warunków 6 i 7 wynika, że równoważność obiektów O_2 i O_4 otrzymujemy wtedy, gdy spełniona jest równość $\lambda_2 = 7 - 13\lambda_1$. (59)

Biorąc pod uwagę, że $0 \leq \lambda_1 \leq 1$, $0 \leq \lambda_2 \leq 1$ oraz $\lambda_1 + \lambda_2 \leq 1$ można wyznaczyć dla różnych parametrów λ_1 i λ_2 wartości składowych wektora $f(p, w_\lambda^4)$ odpowiadające poszczególnym obiektom. Znając wartości tych składowych można określić uporządkowanie obiektów $O_1, \dots, O_4$ odpowiadające przyjętym wartościom λ_1 i λ_2 . W tabeli 1 podano przykładowe wartości λ_1 i λ_2 oraz odpowiadające im wartości składowych dla poszczególnych wektorów, jak również wyznaczone uporządkowania.

Tabela 1

Zależność uporządkowania wynikowego od wartości parametrów λ_1 i λ_2

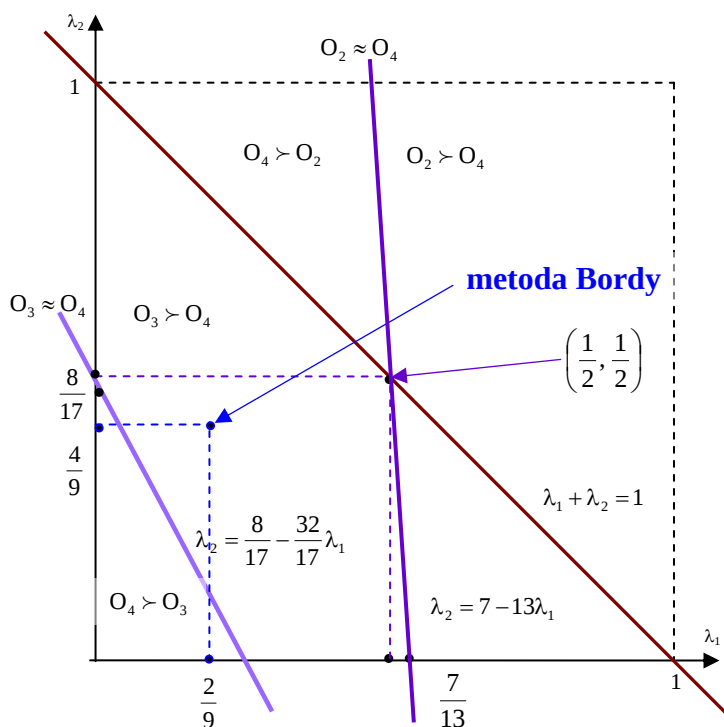
Obiekt		O_1	O_2	O_3	O_4	Uporządkowanie
Składowa wektora $f(p, w_\lambda^4)$		$\frac{34 + 32\lambda_1 + 17\lambda_2}{102}$	$\frac{16 - 10\lambda_2 - 4\lambda_1}{102}$	$\frac{22 + 5\lambda_2 + 2\lambda_1}{102}$	$\frac{30 - 30\lambda_1 - 12\lambda_2}{102}$	
1		2	3	4	5	6
$\lambda_1 = 0$						
$\lambda_2 =$	0	$\frac{34}{102}$	$\frac{16}{102}$	$\frac{22}{102}$	$\frac{30}{102}$	$O_1 \succ O_4 \succ O_3 \succ O_2$
	1	$\frac{51}{102}$	$\frac{6}{102}$	$\frac{27}{102}$	$\frac{18}{102}$	$O_1 \succ O_3 \succ O_4 \succ O_2$

cd. tabeli 1

1	2	3	4	5	6
$\lambda_1 = \frac{1}{4}$					
$\lambda_2 =$	0	$\frac{84}{2 \cdot 102}$	$\frac{30}{2 \cdot 102}$	$\frac{45}{2 \cdot 102}$	$O_1 \succ O_3 \approx O_4 \succ O_2$
	$\frac{3}{4}$	$\frac{219}{4 \cdot 102}$	$\frac{54}{4 \cdot 102}$	$\frac{105}{4 \cdot 102}$	$O_1 \succ O_3 \succ O_4 \approx O_2$
$\lambda_1 = \frac{5}{11}$					
$\lambda_2 =$	0	$\frac{534}{11 \cdot 102}$	$\frac{156}{11 \cdot 102}$	$\frac{252}{11 \cdot 102}$	$O_1 \succ O_3 \succ O_4 \succ O_2$
	$\frac{6}{11}$	$\frac{636}{11 \cdot 102}$	$\frac{96}{11 \cdot 102}$	$\frac{282}{11 \cdot 102}$	$O_1 \succ O_3 \succ O_4 \succ O_2$
$\lambda_1 = \frac{6}{13}$					
$\lambda_2 =$	0	$\frac{634}{13 \cdot 102}$	$\frac{184}{13 \cdot 102}$	$\frac{298}{13 \cdot 102}$	$O_1 \succ O_3 \succ O_4 \succ O_2$
	$\frac{7}{13}$	$\frac{753}{13 \cdot 102}$	$\frac{114}{13 \cdot 102}$	$\frac{327}{13 \cdot 102}$	$O_1 \succ O_3 \succ O_4 \succ O_2$
$\lambda_1 = \frac{7}{13}$					
$\lambda_2 =$	0	$\frac{666}{13 \cdot 102}$	$\frac{180}{13 \cdot 102}$	$\frac{300}{13 \cdot 102}$	$O_1 \succ O_3 \succ O_4 \approx O_2$
	$\frac{6}{13}$	$\frac{768}{13 \cdot 102}$	$\frac{120}{13 \cdot 102}$	$\frac{330}{13 \cdot 102}$	$O_1 \succ O_3 \succ O_2 \succ O_4$
$\lambda_1 = \frac{6}{11}$					
$\lambda_2 =$	0	$\frac{566}{11 \cdot 102}$	$\frac{152}{11 \cdot 102}$	$\frac{254}{11 \cdot 102}$	$O_1 \succ O_3 \succ O_2 \succ O_4$
	$\frac{5}{11}$	$\frac{651}{11 \cdot 102}$	$\frac{102}{11 \cdot 102}$	$\frac{279}{11 \cdot 102}$	$O_1 \succ O_3 \succ O_2 \succ O_4$
$\lambda_1 = 3/4$					
$\lambda_2 =$	0	$\frac{232}{4 \cdot 102}$	$\frac{52}{4 \cdot 102}$	$\frac{94}{4 \cdot 102}$	$O_1 \succ O_3 \succ O_2 \succ O_4$
	$\frac{1}{4}$	$\frac{249}{4 \cdot 102}$	$\frac{42}{4 \cdot 102}$	$\frac{99}{4 \cdot 102}$	$O_1 \succ O_3 \succ O_2 \succ O_4$

Warunki 5 oraz 8 podają zależności, które muszą być spełnione, aby miała miejsce równoważność obiektów O_3 i O_4 oraz obiektów O_2 i O_4 . Proste wyznaczające te zależności przedstawiono na rys. 2.

Z warunków 1, 2 i 3 wynika, że bez względu na wartości parametrów λ_1 i λ_2 zawsze jest spełniona zależność $O_1 \succ O_2$, $O_1 \succ O_3$, $O_1 \succ O_4$. Na rysunku 2 wskazano obszary, dla których $O_2 \succ O_4$ oraz $O_4 \succ O_2$ a także $O_3 \succ O_4$ i $O_4 \succ O_3$.



Rys. 2. Obszary odpowiadające poszczególnym uporządkowaniom

Przykład 3

Jeżeli do uporządkowań $\{P^1, P^2, P^3, P^4\}$ zastosujemy metodę Bordy, to wskaźniki Bordy dla poszczególnych obiektów będą następujące

$$WB_1 = 3 \cdot 11 + 2 \cdot 6 + 1 \cdot 0 = 45 \quad (60)$$

$$WB_2 = 3 \cdot 2 + 2 \cdot 0 + 1 \cdot 6 = 12 \quad (61)$$

$$WB_3 = 3 \cdot 4 + 2 \cdot 5 + 1 \cdot 2 = 24 \quad (62)$$

$$WB_4 = 3 \cdot 0 + 2 \cdot 6 + 1 \cdot 9 = 21 \quad (63)$$

Zatem dla uporządkowań $\{P^1, P^2, P^3, P^4\}$ ocena grupowa według metody Bordy ma postać

$$O_1 \succ O_3 \succ O_4 \succ O_2 \quad (64)$$

Aby wyznaczyć wartości parametrów λ_1 i λ_2 odpowiadających metodzie Bordy, należy zauważyć, że wektor $\mathbf{w}_\lambda^4$ ma postać (4) $\mathbf{w}^{B4} = \left(\frac{3}{4}, \frac{2}{4}, \frac{1}{4}, 0\right)$.

W przykładzie 2 pokazano, że znormalizowany wektor Bordy może mieć postać (9) $\bar{\mathbf{w}}^{B4} = \left(\frac{5}{9}, \frac{3}{9}, \frac{1}{9}, 0\right)$. Stosując opisane podejście do wektora $\bar{\mathbf{w}}^{B4}$ oraz rozpatrywanych uporządkowań $\{P^1, P^2, P^3, P^4\}$ otrzymamy

$$f(\mathbf{p}, \bar{\mathbf{w}}^{B4}) = \frac{5}{17} \left(\frac{5}{9}, 0, \frac{3}{9}, \frac{1}{9}\right) + \frac{2}{17} \left(\frac{3}{9}, \frac{5}{9}, \frac{1}{9}, 0\right) + \frac{4}{17} \left(\frac{3}{9}, 0, \frac{5}{9}, \frac{1}{9}\right) + \frac{6}{17} \left(\frac{5}{9}, \frac{1}{9}, 0, \frac{3}{9}\right) \quad (65)$$

Składowe tego wektora są następujące

$$(i=1) \frac{5}{17} \cdot \frac{5}{9} + \frac{2}{17} \cdot \frac{3}{9} + \frac{4}{17} \cdot \frac{3}{9} + \frac{6}{17} \cdot \frac{5}{9} = \frac{25+6+12+30}{153} = \frac{73}{153} \quad (66)$$

$$(i=2) \frac{5}{17} \cdot 0 + \frac{2}{17} \cdot \frac{5}{9} + \frac{4}{17} \cdot 0 + \frac{6}{17} \cdot \frac{1}{9} = \frac{10+6}{153} = \frac{16}{153} \quad (67)$$

$$(i=3) \frac{5}{17} \cdot \frac{3}{9} + \frac{2}{17} \cdot \frac{1}{9} + \frac{4}{17} \cdot \frac{5}{9} + \frac{6}{17} \cdot 0 = \frac{15+2+20}{153} = \frac{37}{153} \quad (68)$$

$$(i=4) \frac{5}{17} \cdot \frac{1}{9} + \frac{2}{17} \cdot 0 + \frac{4}{17} \cdot \frac{1}{9} + \frac{6}{17} \cdot \frac{3}{9} = \frac{5+4+18}{153} = \frac{27}{153} \quad (69)$$

Zatem uporządkowanie wyznaczone przez $f(\mathbf{p}, \bar{\mathbf{w}}^{B4})$ ma postać

$$O_1 \succ O_3 \succ O_4 \succ O_2. \quad (70)$$

Założmy, że wektor $\bar{\mathbf{w}}^{B4}$ zapiszemy jako sumę wektorów $\mathbf{E}_1^4, \mathbf{E}_2^4, \mathbf{E}_3^4$.

$$\text{Mamy } \left(\frac{5}{9}, \frac{3}{9}, \frac{1}{9}, 0\right) = \lambda_1(1, 0, 0, 0) + \lambda_2\left(\frac{1}{2}, \frac{1}{2}, 0, 0\right) + \lambda_3\left(\frac{1}{3}, \frac{1}{3}, \frac{1}{3}, 0\right) \quad (71)$$

Uwzględniając $\lambda_1 + \lambda_2 + \lambda_3 = 1$ łatwo wykazać, że $\lambda_1 = \frac{2}{9}, \lambda_2 = \frac{4}{9}, \lambda_3 = \frac{3}{9}$.

Na rys. 2 zaznaczono punkt odpowiadający tym wartościom λ_1 i λ_2 . Dla tego punktu mamy $O_3 \succ O_4$ oraz $O_4 \succ O_2$.

Zatem biorąc pod uwagę, że $\forall \lambda_1, \lambda_2$ mamy $O_1 \succ O_3$ można stwierdzić, iż $O_1 \succ O_3 \succ O_4 \succ O_2$. Wynik ten jest zgodny z uzyskanym poprzednio.

Uwagi końcowe

Przedstawione przez Saari'ego geometryczne podejście do pozycyjnych metod wyznaczania oceny grupowej stwarza możliwości wyjaśnienia różnic w ocenach uzyskanych w wyniku zastosowania różnych metod oraz wytłumaczenia występujących paradoksów. Podejście to pozwala również na uwzględnienie możliwości występowania obiektów równoważnych zarówno w ocenach ekspertów, jak i w ocenie grupowej. Umożliwia ono próbę odpowiedzi na pytanie, na ile uzyskane za pomocą różnych metod wyniki odzwierciedlają opinie ekspertów, a w jakim stopniu zależą od zastosowanej metody.

W pracy przeanalizowano przypadek czterech obiektów oraz podano jego interpretację geometryczną.

Literatura

1. Baharad E., Nitzan S. (2003). The Borda Rule, Condorcet Consistency and Condorcet Stability. *Economic Theory*, 22.
2. Baharad E., Nitzan S. (2006). On the Selection of the Same Winner by all Scoring Rules. *Social Choice and Welfare*, 26.
3. Bury H., Wagner D. (2008). Group Judgement With Ties. Distance-Based Methods. In: *New Approaches in Automation and Robotics*. I-Tech Education and Publishing. Ed. H. Aschemann. Vienna, Austria.
4. Saari D.G. (1992). Millions of Election Outcomes from a Single Profile. *Social Choice and Welfare*, 9.
5. Saari D.G. (1995). *Basic Geometry of Voting*. Springer Verlag, Heidelberg.
6. Saari D.G. (2008a). Mathematics and Voting. *Notices of the AMS*, 55, 4.
7. Saari D.G. (2008b). Complexity and the Geometry of Voting. *Mathematical and Computer Modelling*, 48.

Robert Golański

Szkoła Główna Handlowa w Warszawie

ZNACZENIE SENATU W POLSKIM SYSTEMIE LEGISLACYJNYM – PRZESTRZENNY MODEL TEORIOGROWY

Wstęp

W polskim systemie politycznym po przemianach ustrojowych z 1989 roku za przyjmowanie ustaw odpowiedzialne są dwie izby parlamentu, Sejm i Senat. Izby te nie mają charakteru symetrycznego – inaczej niż np. w Stanach Zjednoczonych, w Polsce do przyjęcia ustawy nie jest konieczna zgoda Senatu, sprzeciw izby wyższej może zostać odrzucony przez Sejm. Różnią się one również od 1991 systemem wyborczym.

Różnice w systemie wyborczym przekładają się na różnice w składzie obu izb. 460 posłów na Sejm wybieranych jest w wyborach proporcjonalnych z 5% progiem wyborczym, co w praktyce oznacza ukształtowanie się systemu wielopartyjnego (w Polsce próg wyborczy od kilkunastu lat przekracza 4-6 ugrupowań), w którym żadna partia nie uzyskała dotychczas w wyborach poparcia wystarczającego do uzyskania w Sejmie bezwzględnej większości głosów i umożliwiającego samodzielne tworzenie rządu. 100 senatorów jest natomiast wybieranych w wyborach większościowych – w okręgach wielomandatowych mandaty przypadają kandydatom, którzy uzyskali największą liczbę głosów, przy czym nie ma wymogu, by uzyskane poparcie przekraczało bezwzględną większość ważnie oddanych głosów. System taki prowadzi, zgodnie z obserwacją Duvergera, do wykształcenia się w Senacie systemu dwupartyjnego, i rzeczywiście taka prawidłowość jest obserwowana – oprócz dwóch największych partii pozostałe ugrupowania nie wprowadzają do Senatu reprezentantów lub wprowadzają pojedynczych senatorów, a partia dominująca w Sejmie w Senacie zdobywa bezwzględną większość głosów.

Ścieżka legislacyjna w polskim systemie wygląda następująco: grupy posłów lub rząd (wybierany przez Sejm większością głosów) mogą zgłaszać projekty ustaw^{*}. Po dyskusjach i zgłaszaniu poprawek Sejm przyjmuje projekt

^{*} Inicjatywa ustawodawcza przysługuje również prezydentowi i grupom obywateli, ale przypadki składania ustaw w takim trybie są rzadkie, a ustawy zgłaszane przez podmioty nie mające reprezentacji w Sejmie przyjmowane są relatywnie rzadziej.

ustawy zwykłą większością głosów (więcej głosów za niż przeciw w obecności co najmniej 230 posłów). Uchwalony przez Sejm projekt trafia do Senatu, który może przyjąć projekt bez poprawek, przyjąć poprawki lub odrzucić projekt w całości. Jeśli Senat wprowadzi poprawki lub odrzuci projekt, to Sejm ponownie głosuje nad projektem, odrzucając poprawki Senatu bezwzględną większością głosów. Tak uchwalony projekt trafia do podpisu przez prezydenta, który może projekt podpisać lub odesłać do Sejmu do ponownego rozpatrzenia*. Ewentualne weto prezydenta odrzucane jest przez Sejm większością 3/5 głosów.

Postulat zniesienia Senatu podnoszony jest lub był przez liczące się na polskiej scenie politycznej partie (np. SLD w 2001 czy PO w 2007), w mediach można również spotkać się z opinią, iż fakt, że do odrzucenia weta Senatu wystarczy bezwzględna większość głosów oznacza, że obecność Senatu wydłuża jedynie ścieżkę legislacyjną, ale nie wpływa na to, jakie projekty zostaną przez parlament przyjęte. W niniejszym opracowaniu zbadamy, czy korzystając z prostego modelu przestrzennego możemy pokazać, że istnienie dodatkowego etapu, w którym Senat wyraża swoją opinię, a Sejm się do niej ustosunkowuje (nawet jeżeli do odrzucenia weta Senatu wykorzystuje taką samą większość jak przyjęcie przez Sejm pierwotnego projektu) może rzeczywiście zmienić ostatecznie wybrany przez parlament projekt.

1. Model

Rozważmy grę o następującej strukturze:

1. Zbiór graczy N jest czteroelementowy, $N = \{A, B, C, D\}$. Każdy gracz $i \in N$ charakteryzowany jest przez dwuwymiarowy wektor wag $(x_i; y_i)$, oznaczający liczbę posłów reprezentujących danego gracza odpowiednio w Sejmie i Senacie oraz przez (wspólną dla obu izb) pozycję $p_i = (p_i^1; p_i^2)$ zajmowaną przez danego gracza w dwuwymiarowej przestrzeni. Zakładamy, że gracze są partiami homogenicznymi (wszyscy przedstawiciele danej partii zajmują tę samą pozycję p_i), a użyteczność danego gracza jest dana przez euklidesową odległość od punktu p_i (czyli gracz z dwóch punktów $m = (m^1; m^2)$, $n = (n^1; n^2)$ za lepszy uważa ten, który jest położony bliżej punktu p_i , czyli $m \succsim_i n$ wtedy i tylko wtedy, gdy $(m^1 - p_i^1)^2 + (m^2 - p_i^2)^2 \leq (n^1 - p_i^1)^2 + (n^2 - p_i^2)^2$. Punkt p_i jest oczywiście najbardziej preferowanym punktem gracza i .

* Prezydent może również przekazać ustawę do Trybunału Konstytucyjnego z prośbą o zbadanie konstytucyjności, ale jeśli projekt zostanie w takim trybie uznany za konstytucyjny, to nie może on odmówić podpisania ustawy ani przekazać jej z powrotem do Sejmu.

2. Gracze czerpią użyteczność wyłącznie z tego, jaki punkt został ostatecznie przyjęty w wyniku rozegrania gry jako realizowana przez państwo pozycja. W szczególności gracze nie czerpią żadnej użyteczności bezpośrednio ze sposobu w jaki dokonują oni wyborów. Gracze będą podczas gry zachowywać się strategicznie, podejmując decyzje w taki sposób, aby ostatecznie uzyskać jak najkorzystniejszy wynik z punktu widzenia danego gracza.

3. Istnieje pewien punkt $s = (s^1; s^2)$, odpowiadający status quo. Założymy dla uproszczenia, że dla każdego gracza $i \in N$, s jest punktem ściśle gorszym od każdego punktu, który może zostać przyjęty podczas gry.

4. Gra jest rozgrywana z kompletną informacją, struktura gry, w tym również preferencje poszczególnych graczy, jest wspólną wiedzą wszystkich graczy.

Do rozwiązania przedstawionych gier przyjmiemy jako koncepcję równowagi zaproponowane przez Seltena pojęcie równowagi doskonałej w podgrach, przy czym dla poszczególnych etapów gry – w których decyzje są przez graczy podejmowane równocześnie – ograniczymy się do poszukiwania równowag Nasha korzystając z zaproponowanego przez Moulina pojęcia rozwiązywalności przez dominację*.

Aby pokazać, że Senat w pewnych sytuacjach może pełnić kluczową rolę w procesie legislacyjnym rozważymy dwie gry. W jednej (model 1) rozważymy model jednoizbowy, w którym decyzja podjęta przez Sejm jest ostatecznym wynikiem gry. W drugiej (model 2) gra zostanie zmodyfikowana przez wprowadzenie Senatu, który będzie mógł dokonać poprawek w przyjętym przez Sejm projekcie. Obie gry rozpoczynają się w taki sam sposób. Każdy z graczy może zgłosić propozycję realizowanej polityki, reprezentowaną przez punkt w przestrzeni R^2 . Przyjmiemy, że każda partia może wyłącznie powstrzymać się od zgłoszenia projektu lub zgłosić jako projekt swoją własną pozycję p_i^{**} . Jeśli żaden z graczy nie zgłosił żadnej propozycji, wynikiem gry jest punkt s (status quo). Jeśli projekt został zgłoszony tylko przez jednego gracza, to projekt ten

* Gra jest rozwiązywalna przez dominację, jeśli zastosowanie procedury iteracji eliminacji strategii słabo zdominowanych w taki sposób, że w każdej iteracji usuwane są wszystkie strategie słabo zdominowane każdego gracza, prowadzi ostatecznie do zredukowania gry do postaci, w której każdy gracz jest obojętny pomiędzy każdą parą dowolnych niewykreślonych ostatecznie wyników. Pozwala to uniknąć nieintuicyjnych równowag Nasha, w których gracze mogliby zgłaszać lub głosować na projekty uznawane przez nich za najgorsze, jeśli biorąc zachowanie pozostałych graczy za dane ich pojedynczy głos nie wpływa na wynik.

** Gracze nie są przez wyborców rozliczani ze sposobu, w jaki głosują (i mogą w związku z tym głosować w sposób strategiczny na projekty odległe od ich optymalnego punktu), ale zakładamy, że zgłoszenie przez nich projektu innego od zajmowanej w rzeczywistości pozycji wiązałoby się z silnie negatywną oceną ze strony wyborców, z których opinią gracze się liczą. Zauważmy jednak, że powyższe ograniczenie nie jest bardzo restrykcyjne w tym sensie, że nawet jeśli ograniczymy liczbę projektów, które dana partia może zgłosić, do jednego, to nadal możemy pokazać, że nawet przy tak silnym ograniczeniu wprowadzeniu do gry Senatu może w zasadniczy sposób zmienić wynik gry.

staje się pozycją wybraną przez Sejm. Jeśli zostały zgłoszone dwa projekty, to Sejm w głosowaniu większością zwykłą wybiera jeden z nich^{*}. Jeśli zostały zgłoszone więcej niż dwa projekty, to procedura wyboru wyniku odbywa się przez zastosowanie głosowania przez binarną agendę – po otrzymaniu listy zgłoszonych projektów, marszałek Sejmu ustala kolejność, w której projekty te będą ze sobą kolejno parami porównywane; kolejność (KLM) oznaczałaby np., że w pierwszej kolejności w izbie przeprowadzane jest porównanie pomiędzy K a L, a zwycięzca w tym głosowaniu jest następnie porównywany z M. Zwycięzca ostatniego głosowania jest wybierany jako pozycja zajęta przez Sejm. Gra w modelu 1 kończy się w momencie, w którym Sejm przyjął jakąś pozycję (lub wyborem status quo jeśli żaden gracz nie zgłosił żadnego projektu), a pozycja wybrana przez Sejm staje się wynikiem gry.

W modelu drugim projekt przyjęty przez Sejm nie staje się wynikiem gry, ale trafia do Senatu. W Senacie każda partia może zgłosić poprawkę (odpowiadającą podobnie jak w Sejmie zajmowanej przez daną partię pozycji). Jeśli nie została zgłoszona żadna poprawka (lub jeśli zgłoszona została wyłącznie poprawka identyczna z projektem zgłoszonym przez Sejm), to projekt Sejmu staje się wynikiem gry i gra się kończy. Jeśli natomiast w Senacie zostały zgłoszone poprawki, to marszałek Senatu ustala kolejność, w jakiej spośród zgłoszonych poprawek w procedurze głosowania przez agendę binarną zostanie wyłoniony projekt Senatu. Jeśli projekt wybrany przez Senat jest różny od projektu przedstawionego przez Sejm, to w Sejmie przeprowadzane jest jeszcze jedno głosowanie, w którym Sejm dokonuje porównania między oboma projektami i wybiera jeden z nich w głosowaniu większością bezwzględną^{**} – wynikiem gry jest w takiej sytuacji projekt wybrany w tym głosowaniu. Założymy, że w modelu nie ma prezydenta, a zatem przyjęcie projektu ustawy zależy wyłącznie od decyzji parlamentarzystów.

Zgodnie z prawidłowościami obserwowanymi w rzeczywistości na polskiej scenie politycznej założymy w powyższym modelu, że jeden z czterech graczy (mianowicie gracz A) jest graczem dominującym. Oznaczać to będzie dwie zależności:

^{*} W przypadku remisu założymy, że marszałek Sejmu ma możliwość wyboru, który z dwóch projektów zostanie wybrany. W przykładzie rozważanym w niniejszym opracowaniu nie będziemy mieli do czynienia z tego typu sytuacją

^{**} Ponieważ w przedstawionym przykładzie żaden z graczy nie wykazuje obojętności pomiędzy żadnymi dwoma możliwymi projektami, to nikt nie będzie miał bodźca do tego, żeby w równowadze wstrzymać się od głosowania w sytuacji, gdy mogłoby to wpłynąć na ostateczny wynik – założenie o tym, że poprawki Senatu są odrzucane wyższą (bezwzględną) większością niż projekt przyjmowany w pierwszym kroku przez Sejm (który przyjmowany jest większością zwykłą) nie ma wpływu na uzyskane wyniki – podobny rezultat uzyskalibyśmy nawet gdybyśmy założyli, że pozycja Senatu jest słabsza niż w polskim systemie ustrojowym i do odrzucenia jego weta nie potrzeba podwyższonej większości, a większość bezwzględna stosowana jest na obu etapach gry.

1. Gracz A kontroluje Senat w tym sensie, że y_A jest co najmniej równe minimalnej liczbie głosów niezbędnej do podjęcia decyzji w Senacie (czyli jest równe co najmniej 51) i partia A kontroluje stanowisko marszałka Senatu. Koalicjami wygrywającymi w Senacie będą zatem takie i tylko koalicje, w skład których wchodzi gracz A, a gracz A staje się w Senacie dyktatorem w sensie Arrowa – zagregowane preferencje grupowe Senatu jako izby pokrywają się z preferencjami gracza A, a ponadto z preferencjami gracza A pokrywają się również preferencje marszałka Senatu.

2. W Sejmie dominująca pozycja gracza A oznacza, że gracz A dysponuje liczbą głosów wystarczającą do tego, by utworzyć koalicję wygrywającą z każdym z pozostałych trzech graczy (czyli dla każdego $i \in \{B, C, D\}$ zachodzi $x_A + x_i > 230$). Pozostali trzej gracze dysponują jednak liczbą głosów wystarczającą do podjęcia decyzji nawet bez zgody gracza A (czyli $x_B + x_C + x_D > 230$). Rozpatrzmy w niniejszym artykule zarówno przypadek, w którym preferencje marszałka Sejmu odpowiadają preferencjom gracza A, jak i przypadek gdy odpowiadają preferencjom jednego z pozostałych graczy.

2. Przypadek jednowymiarowy

Rozważmy najpierw sytuację, w której wybierane projekty znajdują się na wspólnej prostej – załóżmy, że mamy $p_A^2 = p_B^2 = p_C^2 = p_D^2 = s^2$, i podobnie dla każdego zgłoszonego projektu (m^1, m^2) musi zachodzić $m^2 = s^2$. Uchylmy natomiast założenie, że gracze mogą zgłaszać wyłącznie projekty odpowiadające dokładnie zajmowanym przez nich pozycjom p_i^* . Zauważmy jednak, że w przypadku, kiedy przestrzeń polityczna jest jednowymiarowa, to przy przyjętych założeniach dotyczących liczby posłów, którymi dysponują poszczególni gracze (w szczególności przy założeniu, że nie istnieje koalicja dysponująca dokładnie 230 głosami), w grze istnieje dokładnie jeden punkt odpowiadający pozycji medianowego posła – a zatem pozycja medianowego posła odpowiadać będzie również pozycji jednej z partii, p_i (gdyż założyliśmy, że partie są homogeniczne). Oznaczmy taki punkt przez p . Oczywiście jest jednak, że punkt p jest w Sejmie w takiej grze zwycięzcą w sensie Condorceta – porównując punkt p z jakimkolwiek innym punktem, który mógłby zostać zgłoszony, zawsze będzie istniała bezwzględna większość graczy, którzy uważają punkt p za preferowany. Rozważmy model 1 (jednoizbowy). Jeśli istnieje zwycięzca w sensie Condorceta, to niezależnie od tego, ile projektów zostanie zgłoszonych i w jakiej kolejności będą one ze sobą porównywane, to jeśli wszyscy gracze – zgodnie

* Uchylenie tego założenia nie wpływa na równowagę takiej gry – nawet jeśli gracze mogą zgłaszać dowolne projekty, to jak pokażemy w równowadze wybrany zostanie punkt odpowiadający optymalnemu punktowi jednego z graczy.

z przyjętymi założeniami – głosują strategicznie i zainteresowani są wyłącznie ostatecznym rezultatem głosowania, to jeśli tylko zwycięzca w sensie Condorceta został zgłoszony jako jeden z projektów, to przy zastosowaniu głosowania przez agendę binarną zostanie on ostatecznie wybrany. Oczywiście jest jednak, że projekt taki zostanie zgłoszony – ponieważ punkt p odpowiada optymalnemu punktowi jednego z graczy, to gracz ten, mogąc w równowadze zagwarantować sobie, że punkt p zostanie wybrany, jeśli tylko zostanie zgłoszony, zgłosi oczywiście punkt p jako jeden z projektów. W przypadku zatem, kiedy wszystkie możliwe projekty można umieścić w jednym wymiarze, a decyzja jest podejmowana wyłącznie przez Sejm, to niezależnie od tego, jak wygląda struktura głosów w Sejmie oraz która partia kontroluje stanowisko marszałka, wybrany zostanie zawsze ten sam punkt p.

Zauważmy jednak, że przyjęty propozycja będzie również równa p jeżeli rozpatrzymy model 2 (dwuizbowy). Gra składa się w modelu 2 z trzech etapów, w których kolejno najpierw Sejm przyjmuje swoją propozycję, potem Senat może przyjąć poprawkę (przedstawiając kontrpropozycję) i wreszcie Sejm dokonuje ostatecznego wyboru między propozycją Senatu a oryginalną propozycją Sejmu. Jeżeli jednak Sejm przyjął w pierwszym etapie projekt p, to – ponieważ jest on zwycięzcą w sensie Condorceta – niezależnie od tego czy Senat w drugim etapie przedstawi jakąś kontrpropozycję, to zostanie ona w trzecim etapie odrzucona na korzyść p.

Podobnie w sytuacji, kiedy Sejm co prawda nie wybrał propozycji p, ale p zostało zaproponowane przez Senat – ostatecznie w porównaniu w trzecim etapie Sejm wybierze p, gdyż p jest zwycięzcą w sensie Condorceta.

Ostatnim przypadkiem jest sytuacja, kiedy ani Sejm, ani Senat nie proponowały p jako swojego projektu. To jednak nie może się zdarzyć w równowadze, gdyż w takiej sytuacji ostatecznym rezultatem byłby projekt p' od p różny. Gracze są jednak strategiczni i zainteresowani wyłącznie ostatecznym rezultatem gry. Skoro zatem p jest w Sejmie zwycięzcą w sensie Condorceta, to istnieje w Sejmie większość graczy, która wolałaby p od p'. Ponieważ jednak zgłoszenie w Sejmie projektu p oznacza, że projekt ten zostanie ostatecznie przyjęty, to w równowadze musi istnieć gracz, który zgłosi w Sejmie taki projekt i w efekcie zostanie on przyjęty przez Sejm, nie może zatem dojść w równowadze do sytuacji, gdzie p nie zostało wybrane przez żadną z izb.

W przypadku modelu jednowymiarowego wprowadzenie do modelu Senatu nie wpływa zatem na ostatecznie przyjęty projekt. Nie zależy to od liczby wymiarów przestrzeni, z której pochodzą projekty – identyczny wynik zostanie uzyskany niezależnie od liczby wymiarów, jeśli tylko wśród możliwych do zgłoszenia projektów istnieje zwycięzca w sensie Condorceta.

3. Przypadek dwuwymiarowy

Przy założeniu, że projekty pochodzą z przestrzeni dwuwymiarowej, jeśli gracze mogliby składać dowolne propozycje, należy się spodziewać, że zwycięzca w sensie Condorceta nie będzie istniał – warunki Plotta określające konieczne warunki do tego, by zwycięzca w sensie Condorceta istniał są silnie restrykcyjne. Jeśli jednak ograniczymy możliwość zgłaszania projektów wyłącznie do punktów p_i odpowiadających optymalnym punktom poszczególnych partii, to nietrudno skonstruować przykład, w którym jeden z takich punktów jest zwycięzcą w sensie Condorceta. Rozważmy np. sytuację, w której $p_A = (0, 4)$, $p_B = (-4, 0)$, $p_C = (1, -1)$ i $p_D = (3, -2)$. Przy takim rozmieszczeniu pozycji poszczególnych graczy punkt p_C jest zwycięzcą w sensie Condorceta – gracz C woli punkt p_C od każdego innego punktu, ale gracz A woli p_C zarówno od p_B , jak od p_D (a koalicja graczy A i C jest koalicją wygrywającą), natomiast gracze B i D wolą punkt p_C od p_A (i podobnie, koalicja graczy B, C i D dysponuje bezwzględną większością głosów). Gdyby zatem pozycje graczy znajdowałyby się w powyżej opisanych punktach, to zwycięzca w sensie Condorceta istnieje, a w takim razie zostanie on wybrany niezależnie od tego, czy w modelu znajduje się dodatkowy etap w postaci poprawek Senatu czy nie.

Możemy zwrócić jednak uwagę na dwie kwestie – po pierwsze, powyższe cztery punkty nie spełniają warunków Plotta, a zatem gdyby gracze mogli składać dowolne propozycje, to istniałyby punkty preferowane przez większość od punktu p_C (np. gracze A i B wolą punkt $(0, 0)$ od punktu p_C , i gracze ci tworzą koalicję wygrywającą, a zatem nie dla każdego punktu istnieje większościowa koalicja, która uważa punkt p_C za lepszy). Po drugie jednak, nawet jeśli ograniczymy możliwość zgłaszania przez daną partię projektów do pozycji rzeczywiście zajmowanych przez daną partię, p_i , to również wtedy może się zdarzyć, że żaden z takich punktów nie będzie zwycięzcą w sensie Condorceta.

Rozważmy następujący przykład. Niech $p_A = (0, 3)$, $p_B = (-4, -1)$, $p_C = (3, -1)$ i $p_D = (1, -3)$. Przy tak określonych pozycjach zajmowanych przez poszczególnych graczy, ich preferencje są zatem następujące:

dla gracza A: $p_A \succ_A p_C \succ_A p_B \succ_A p_D$

dla gracza B: $p_B \succ_B p_D \succ_B p_A \succ_B p_C$

dla gracza C: $p_C \succ_C p_D \succ_C p_A \succ_C p_B$

dla gracza D: $p_D \succ_D p_C \succ_D p_B \succ_D p_A$

Zagregowanie preferencji za pomocą głosowania większością bezwzględną (lub zwykłą – żaden z graczy nie jest tu obojętny między żadnymi dwoma dopuszczalnymi punktami i w takim razie obie metody dadzą takie

same preferencje grupowe) przy poczynionym założeniu, że gracz A może stworzyć koalicję wygrywającą z każdym z pozostałych graczy, ale sam ma mniej niż 230 głosów, oznacza że zagregowane preferencje Sejmu są następujące: $p_A \succ p_B$, $p_A \succ p_C$, $p_D \succ p_A$, $p_B \succ p_D$, $p_C \succ p_D$, $p_C \succ p_B$. Gra nie ma zatem zwycięzcy w sensie Condorceta – grupa nie uważa większością głosów za projekt lepszy od każdego innego ani A (D jest lepsze), ani B czy C (A jest lepsze), ani D (C jest lepsze). Oznacza to w szczególności, że gdyby wszystkie cztery projekty zostały zgłoszone, to to, który z nich wygrałby w Sejmie głosowanie przez agendę binarną zależeć będzie od kolejności, w jakiej projekty zostaną pod głosowanie, a zatem zależeć będzie od tego, który gracz kontroluje stanowisko marszałka, który kolejność głosowania będzie ustalał. Decyzja o kolejności poddania wniosków pod głosowanie może doprowadzić do różnych wyników również w dwóch innych przypadkach, dla których również występuje cykliczność preferencji grupowych, mianowicie jeśli zgłoszone zostały wyłącznie projekty p_A , p_B , p_D lub jeśli zgłoszone zostały wyłącznie projekty p_A , p_C , p_D .

Rozważmy model 1 (bez Senatu), w którym stanowisko marszałka Sejmu kontrolowane jest przez gracza A. Gra taka będzie zatem przebiegała w taki sposób, że najpierw każdy gracz $i \in N$ równocześnie z innymi decyduje, czy zgłosić swój optymalny punkt p_i czy nie. Jeśli zostało zgłoszonych więcej niż 2 projekty, marszałek Sejmu wybiera kolejność, w jakiej będą one ze sobą kolejno parami porównywane (zwycięzca danego głosowania porównywany jest następnie z kolejnym projektem w wybranej agendzie), partie zachowują się strategicznie tak, aby uzyskać jako rezultat głosowania jak najkorzystniejszy dla siebie wynik, a projekt wybrany w głosowaniu jest pozycją zajęta przez Sejm.

Zauważmy, że w każdej z trzech sytuacji, w których zwycięzca w sensie Condorceta nie istnieje, istnieje taka kolejność głosowania, w której – jeśli gracze zachowują się optymalnie – wybrany zostanie optymalny punkt gracza A, p_A : jeśli zgłoszone zostały wszystkie projekty, to taką kolejnością jest np. (p_B przeciwko p_D , zwycięzca przeciwko p_A , zwycięzca drugiego głosowania przeciwko p_C), jeśli zgłoszone zostały wyłącznie projekty p_A , p_B , p_D , to taką kolejnością jest (p_A przeciwko p_D , zwycięzca przeciwko p_B), a jeśli zgłoszone zostały wyłącznie projekty p_A , p_C , p_D , to taką kolejnością jest (p_A przeciwko p_D , zwycięzca przeciwko p_C)*. Tabela 1 podsumowuje, który projekt zostanie z takim modelem ostatecznie wybrany w zależności od tego, która partia zdecydowała się zgłosić swój projekt (Z oznacza, że dana partia zgłosiła projekt, a N że nie zgłosiła).

* Zauważmy, że większość graczy woli p_D od p_A . Gracze są jednak strategiczni i wiedzą, że głosowanie na p_D doprowadzi ostatecznie do wyboru p_C , które przez większość graczy uważane jest za gorsze od p_A – gracz B zgłasza zatem strategicznie w pierwszym głosowaniu na p_A , chociaż woli p_D od p_A , aby ostatecznie uzyskać korzystniejszy dla siebie od p_C wynik p_A . Podobnie w przypadku pozostałych sytuacji.

Tabela 1

		Gracz D							
		Z				N			
Gracz C	Z			Gracz B				Gracz B	
				Z	N			Z	N
		Gracz A	Z	A	A	Gracz A	Z	A	A
			N	C	C		N	C	C
	N			Gracz B				Gracz B	
				Z	N			Z	N
		Gracz A	Z	A	D	Gracz A	Z	A	A
			N	B	D		N	B	S

Do rozwiązywania gry stosujemy koncepcję równowagi doskonałej – rezultatem każdej podgry następującej po etapie zgłaszania przez partie projektów jest, w zależności od tego, które projekty zostały zgłoszone, projekt który odczytać możemy z tabeli 1. Gracze będą zatem decydować o tym, czy podjąć decyzję Z czy N kierując się ostatecznym rezultatem gry. Zauważmy jednak, że powyższa gra rozwiązywalna jest przez dominację: gracze A i D wybierając Z uzyskują zawsze wynik nie gorszy niż wybierając N*. Ale w takim razie zgłoszone zostaną projekty p_A i p_D – i jeśli żaden inny projekt nie zostanie zgłoszony, to wygra p_D . Jednak jeśli zgłoszony zostanie jakiś inny projekt, to powstanie cykl w preferencjach grupowych i marszałek Sejmu będzie mógł wnioski poddać pod głosowanie w taki sposób, żeby wygrało p_A . Gracze B i C wolą jednak p_D od p_A , aby zatem uniknąć takiej sytuacji, żaden z nich nie zgłosi własnego projektu. W równowadze takiej gry zostaną zatem zgłoszone wyłącznie p_A i p_D , a zwycięskim projektem zostanie p_D .

Zauważmy, że wygrywa w takim razie projekt uważany za najgorszy przez dominującą partię, która dysponuje największą liczbą głosów oraz kontroluje stanowisko marszałka Sejmu. Dzieje się tak dlatego, że gracze B i C przewidując, że marszałek Sejmu skorzysta ze swojego prawa do ustalenia kolejności głosowania w taki sposób, który spowoduje zwycięstwo projektu uważanego przez nich za gorszy, powstrzymują się strategicznie od zgłaszania własnych propozycji. Przyjrzyjmy się alternatywnej sytuacji, która różni się od powyższej wyłącznie tym, która partia kontroluje stanowisko marszałka. Załóżmy, że partią tą jest gracz C (może to np. być gracz, z którym partia A tworzy koalicję rządową).

* Zauważmy również, że możemy osłabić założenie o tym, że s (status quo) jest uważane za najgorszy punkt przez każdego gracza – w powyższym przykładzie wystarczy, że założymy, że gracz A lub gracz D nie uważa punktu s za punkt lepszy od odpowiednio p_A lub p_D .

Jeżeli preferencje marszałka odpowiadają preferencjom gracza C, to zauważmy, że podobnie jak we wcześniejszym przykładzie w trzech sytuacjach, w których zagregowane preferencje grupowe są cykliczne, marszałek będzie mógł tak dobrać kolejność głosowania, żeby uzyskać wynik dla siebie jak najkorzystniejszy. Jeśli zatem zgłoszone zostały wszystkie projekty lub jeśli zgłoszono projekty p_A , p_C , p_D , to zwycięży projekt p_C , jeśli natomiast zgłoszono projekty p_A , p_B , p_D , to wybrana zostanie kolejność prowadząca do wyboru najlepszego spośród nich wariantu ze względu na preferencje gracza C, czyli projekt p_D . Analogiczna tabela opisująca zwycięskie warianty w każdej podgrze następującej po etapie zgłaszania projektów to tabela 2.

Tabela 2

		Gracz D							
		Z				N			
				Gracz B				Gracz B	
				Z	N			Z	N
Gracz C	Z	Gracz A	Z	C	C	Gracz A	Z	A	A
			N	C	C		N	C	C
				Gracz B				Gracz B	
				Z	N			Z	N
	N	Gracz A	Z	D	D	Gracz A	Z	A	A
			N	B	D		N	B	S

Również ta gra jest rozwiązywalna przez dominację – dla gracza C i D wybór decyzji Z jest teraz nie gorszy od wyboru decyzji D. Ale w takim razie zwycięskim wariantem będzie projekt p_C , niezależnie od tego, czy gracze A i B zgłoszą własne projekty czy nie. Wygrywa zatem projekt p_C .

Zauważmy, że oddanie przez dominującą partię stanowiska marszałka jednej z mniejszych partii spowodowało zatem, że gracz dominujący uzyskał wynik korzystniejszy dla siebie – punkt p_C jest przez gracza A uważany za korzystniejszy od punktu p_D . Powstaje zatem paradoks podobny do paradoksu przewodniczącego Farquharsona – kontrolowanie stanowiska odpowiadającego za ustalanie kolejności głosowania i wpływanie na wynik w przypadku sytuacji kontrowersyjnych może pogorszyć wynik z punktu widzenia preferencji gracza, który takie stanowisko kontroluje.

Pokazaliśmy zatem, że w modelu 1 (jednoizbowym), jeśli gracz A kontroluje stanowisko marszałka, to przy tak rozłożonych punktach optymalnych poszczególnych graczy zwycięskim punktem będzie punkt p_D . Rozważmy następnie model 2 (dwuizbowy), dodając kolejny etap, w którym Senat będzie mógł zaproponować poprawkę do propozycji zgłoszonej przez Sejm, które to

następnie dwie propozycje będą w kolejnym etapie porównane przez Sejm. Zgodnie z przyjętymi założeniami przyjmiemy, że Senat jest kontrolowany przez partię A – preferencje Senatu odpowiadają preferencjom gracza A i gracz A kontroluje stanowisko marszałka Senatu.

Gra jest grą z kompletną informacją i rozwiązujemy ją korzystając z pojęcia równowagi doskonałej, Senat będzie zatem podejmował decyzje przewidując, jak zachowa się Sejm, jeśli Senat zaproponuje poprawki. Zauważmy, że w takim razie optymalny wybór Senatu jest następujący:

- jeśli Sejm przyjął w pierwszym etapie projekt p_A , to jest to projekt maksymalizujący już preferencje Senatu, więc Senat nie zaproponuje żadnej poprawki,
- jeśli Sejm przyjął w pierwszym etapie projekt p_B lub projekt p_C , to Senat zaproponuje projekt p_A – Sejm porównując następnie poprawkę Senatu z własnym oryginalnie przyjętym projektem przyjmie projekt p_A – dla każdego projektu p_B i p_C istnieje bowiem większość, preferująca p_A od oryginalnej propozycji Sejmu, i również Senat uważa p_A za projekt lepszy,
- jeśli wreszcie Sejm przyjął w pierwszym etapie projekt p_D , to Senat nie zaproponuje jako poprawki projektu p_A (bo preferencje Sejmu są $p_D > p_A$ i Sejm taką poprawkę odrzuci), ale może zaproponować jako poprawkę projekt p_C , który z punktu widzenia preferencji partii A (czyli również Senatu) jest projektem lepszym od p_D , a który zostanie wybrany jako lepszy również przez Sejm*.

W takim razie jednak w modelu 2 jedynymi możliwymi wynikami są p_A , p_C lub s (w równowadze s oczywiście nie zostanie wybrane), przy czym p_C zostanie wybrane jeśli pierwotnie Sejm wybrał p_D , natomiast p_A zostanie wybrane w pozostałych przypadkach. Głosowanie na p_D w pierwszym głosowaniu oznacza zatem ostateczny wybór p_C , a głosowanie na jakikolwiek inny projekt oznacza ostateczne zwycięstwo p_A . Gracze dysponujący większością głosów (A i B) wolą jednak p_A od p_C , zatem jedyny sposób, w jaki p_D mogłoby wygrać w pierwszym etapie (co w dalszych etapach gry doprowadziłoby do ostatecznego zwycięstwa p_C) to sytuacja, w której p_D było jedynym zgłoszonym projektem – jeśli zgłoszono jakikolwiek inny projekt, to gracze A i B zagłosują strategicznie na jakikolwiek inny projekt, aby Senat mógł jako poprawkę zaproponować następnie p_A . Tabela z projektami ostatecznie przyjętymi w wyniku gry z modelu 2 to tabela 3.

* Przy przyjętych ograniczeniach dotyczących tego, że partia może zgłosić tylko własny optymalny punkt, p_C musiałoby w Senacie zostać zgłoszone w takim razie przez gracza C – ale gracz C w równowadze będzie chciał taką propozycję zgłosić, jeśli będzie ona miała szanse wygrać, a jeśli taka propozycja zostanie zgłoszona, to partia A zagłosuje w Senacie strategicznie właśnie na tę propozycję.

Tabela 3

		Gracz D							
		Z				N			
Gracz C	Z			Gracz B				Gracz B	
				Z	N			Z	N
		Gracz A	Z	A	A	Gracz A	Z	A	A
			N	A	A		N	A	A
	N			Gracz B				Gracz B	
				Z	N			Z	N
		Gracz A	Z	A	A	Gracz A	Z	A	C
			N	A	A		N	A	S

Ponownie zauważmy, że gra jest rozwiązywalna przez dominację – gracz B woli p_A od p_C i jest w stanie zagwarantować sobie wybór p_A zgłaszając w pierwszym etapie projekt p_B – projekt p_D nie będzie wówczas jedynym zgłoszonym projektem (jeśli w ogóle zostanie zgłoszony), a partia A będzie miała możliwość albo strategicznego zgłoszenia na p_B (aby następnie poprawić w Senacie taki projekt na p_A), albo spowodować wybór p_A w pierwszym etapie, bez konieczności modyfikowania wyniku w dalszych etapach gry.

Podsumowanie

Pokazaliśmy w niniejszym opracowaniu przykłady ilustrujące następujące dwa zjawiska:

1. Jeśli w Sejmie nie ma punktu, który byłby zwycięzcą w sensie Condorceta, to kontrolowanie stanowiska marszałka Sejmu może wpłynąć na wynik gry i na to, który projekt zostanie przyjęty, ale wpływa w sposób, który może być nieintuicyjny – rezygnacja ze stanowiska marszałka może w równowadze polepszyć wynik uzyskiwany przez danego gracza.

2. W sytuacji, kiedy żaden punkt nie jest zwycięzcą w sensie Condorceta, wprowadzenie do modelu Senatu może zmienić wynik gry. Zauważmy również, że jeśli założymy, że rozkład sił w Sejmie odpowiada w przybliżeniu rozkładowi preferencji w społeczeństwie i że dominująca partia ma dostatecznie duże poparcie w stosunku do pozostałych partii (na przykład zakładając, że partie B, C i D mają równe poparcie, a partia A poparcie dwukrotnie większe od każdej z nich z osobna), to mierząc dobrobyt społeczny sumą odległości od przyjętej przez parlament polityki, punkt p_A (przyjęty w modelu dwuizbowym) jest społecznie lepszy według tego kryterium od punktu p_D (przyję-

tego w modelu jednoizbowym). Wprowadzenie do modelu kolejnego etapu polegającego na ocenie przez Senat projektu wybranego przez Sejm może nie tylko zmienić, ale polepszyć ze społecznego punktu widzenia wybrany przez parlament projekt.

Powyższy model może zostać rozszerzony na kilka sposobów, które mogą być punktem wyjścia do dalszych badań:

1. Partie w rzeczywistości nie są heterogeniczne. O ile zasadne wydaje się założenie, że nie każdy punkt w przestrzeni dwuwymiarowej może zostać zgłoszony przez daną partię, to poszczególni posłowie różnią się poglądami i dana partia może w takim razie zgłaszać różne projekty z pewnym sąsiedztwie pozycji odpowiadającej uśrednionej pozycji danej partii czy pozycji jej lidera.

2. Rzeczywista ścieżka legislacyjna w Polsce obejmuje również kolejny – czwarty – etap, w którym prezydent może zwrócić do Sejmu projekt z prośbą o ponowne rozpatrzenie. Jeśli uchylimy założenie, że status quo nie jest uważane przez wszystkie partie za najgorszy punkt (a uchylenie tego założenia wydaje się zasadne – status quo może być pozycją wybraną przez parlament poprzedniej kadencji i może być w takim razie samo pozycją zajmowaną przez jednego z graczy), to dodanie takiego kolejnego etapu również może wpłynąć na to, jaki projekt zostanie ostatecznie wybrany.

3. Krzywe obojętności graczy nie muszą być okręgami – różni gracze mogą w różny sposób oceniać, który z dwóch wymiarów jest dla nich ważniejszy, a zatem preferencje poszczególnych graczy mogą być eliptyczne.

Literatura

1. Condorcet J. markiz de (1785/1976). *Essays on the Application of Mathematics to the Theory of Decision Making*. W: Condorcet: Selected Writings. Red. K.M. Baker. Indianapolis, Bobbs-Merrill.
2. Farquharson R. (1969). *Theory of Voting*. Yale University Press, New Haven.
3. Moulin H. (1979). Dominance Solvable Voting Schemes. *Econometrica*, 47, 1337-1351.
4. Nash J.F. (1950). Equilibrium Points in N-Person Games. *Proceedings of the National Academy of Sciences of the United States of America*, 36, 48-49.
5. Plott Ch. (1967). A Notion of Equilibrium and Its Possibility Under Majority Rule. *American Economic Review*, 57, 787-806.
6. Selten R. (1965). Spieltheoretische Behandlung eines Oligopolmodells mit Nachfragetraegheit. *Zeitschrift fuer die gesamte Staatswissenschaft*, 121, 301-324.

Leszek Klukowski

Instytut Badań Systemowych PAN w Warszawie

OGRANICZANIE RYZYKA KOSZTÓW OBSŁUGI DŁUGU PUBLICZNEGO PRZY WYKORZYSTANIU ROZWIĄZANIA MINIMAKSOWEGO GRY STRATEGICZNEJ

Wprowadzenie

Minimalizacja kosztów obsługi portfela instrumentów dłużnych emitowanych przez Skarb Państwa (bonów i obligacji skarbowych) opiera się na długoterminowych prognozach rynkowych stóp procentowych [4-9]; ich horyzont obejmuje okres życia tych instrumentów, obecnie do 30 lat. Prognozy takie mają często postać wariantową, tj. skończonego zbioru ścieżek przyszłych stóp, przy czym nie dysponujemy zwykle wiarygodnymi ocenami prawdopodobieństw wystąpienia poszczególnych ścieżek. Sprzedaż portfela opartego na błędnej prognozie, tj. określonego na podstawie wariantu różnego od zrealizowanego, powoduje stratę, polegającą na wzroście kosztów obsługi długu. Zadaniem emitenta instrumentów dłużnych jest minimalizacja ryzyka takiej straty, tj. realizacja postulatów „bezpiecznego” działania w obszarze finansów publicznych.

Celem niniejszej pracy jest przedstawienie koncepcji rozwiązania tego problemu, polegającej na zastosowaniu rozwiązania minimaksowego dwuosobowej gry strategicznej, ze skończoną liczbą strategii, o sumie zero. Graczami w tej grze są: emitent instrumentów dłużnych (Skarb Państwa) oraz rynek finansowy (nabywcy instrumentów dłużnych); graczem optymalizującym swoje decyzje jest emitent. Zbiór strategii tej gry jest określony przez zbiór wariantów prognoz, a macierz gry zawiera elementy wyrażające przyrost kosztów obsługi spowodowany sprzedażą portfela odpowiadającego ustalonemu wariantowi prognozy, w przypadku realizacji innego wariantu. Rozwiązanie minimaksowe rozważanej gry ma, w ogólnym przypadku, postać strategii zrandomizowanej; zapewnia ono emitentowi najostrożniejszą decyzję, wyrażającą się przez wartość gry [2; 10].

Rozważany w pracy problem można również sformułować w postaci gry z naturą, uznając, że drugi gracz nie generuje strategii w sposób antagonisticzny. Czynnikiem skłaniającym do zastosowania gry strategicznej jest, poza postulatami „ostrożnego” działania, trudność określenia warunkowych rozkładów prawdopodobieństwa stanów natury.

Wdrożenie proponowanej koncepcji ma szczególne znaczenie w warunkach zmienności sytuacji na rynku finansowym oraz zagrożenia kryzysem w obszarze finansów publicznych.

W praktyce zarządzania długiem stosowane są różnorodne metody ograniczania ryzyka, między innymi oparte na metodologii *cost at risk* i technikach symulacyjnych [1]. Idea przedstawiona w niniejszej pracy poszerza spektrum istniejących metod oraz nie ma wad metodologii *cost at risk*. Do istotnych cech proponowanej koncepcji należy zaliczyć: dokładne odzwierciedlenie problemu decyzyjnego, brak krępujących założeń oraz możliwość realizacji obliczeń przy wykorzystaniu powszechnie używanego oprogramowania, np. systemu Excel. Oryginalnym wkładem metodologicznym jest ponadto sposób określenia elementów macierzy gry, uwzględniający realia sytuacji decyzyjnej emitenta instrumentów dłużnych.

Praca składa się z 5 części; kolejne części zawierają: sformułowanie problemu, koncepcję minimaxowego rozwiązania gry, wraz z definicją macierzy gry, przykład zastosowania oraz podsumowanie.

1. Sformułowanie problemu

Problem minimalizacji ryzyka kosztów obsługi długu publicznego, w warunkach posiadania wielowariantowych prognoz stóp procentowych oraz zestawu emitowanych instrumentów dłużnych, można sformułować w następujący sposób.

Dany jest u -elementowy ($2 \leq u < \infty$) zbiór wariantowych prognoz stóp procentowych oraz zbiór odpowiadających im optymalnych portfeli emitowanych instrumentów dłużnych. Sprzedaż portfela odpowiadającego j -temu wariantowi prognozy, w przypadku wystąpienia wariantu i -tego, powoduje przyrost kosztów obsługi długu (kwoty odsetek) wynoszący ω_{ij} ($1 \leq i, j \leq u$; $\omega_{ij} \geq 0$; $\omega_{ij} = 0$); warianty prognozy są wynikiem sytuacji na rynku finansowym. Macierz $\mathbf{W} = [\omega_{ij}]$ ($i, j = 1, \dots, u$) określa wszystkie możliwe przyrosty kosztów (straty emitenta), wynikające z błędów w wyborze wariantu prognozy.

Należy określić sposób wyboru wariantu prognozy przez emitenta, zapewniający minimalizację ryzyka przyrostu kosztów obsługi długu, dla danej macierzy $\mathbf{W}$, przy założeniu, że emitent i rynek działają jak gracze antagonisticzni w grze strategicznej, tj. podejmują decyzje równocześnie i niezależnie.

Sformułowany powyżej problem można sformalizować przy wykorzystaniu koncepcji dwuosobowej gry strategicznej ze skończoną liczbą strategii, o sumie zero, z zastosowaniem rozwiązania minimaxowego [2; 10].

Elementami gry, stanowiącej model rozważanego problemu decyzyjnego, są: gracze – emitent instrumentów dłużnych oraz rynek (inwestorzy kupujący instrumenty dłużne i otrzymujący odsetki), zbiór strategii graczy (zbiór wariantów prognoz), macierz gry (macierz strat emitenta) $\mathbf{W}$ określająca skutki błędu w wyborze wariantu przez emitenta.

Rozwiązanie optymalne gry z macierzą $\mathbf{W}$ ma postać zrandomizowaną, tj. rozkładu prawdopodobieństwa na zbiorze strategii (wariantów prognoz); w przypadku, gdy macierz ta ma punkt siodłowy, rozkład jest jednopunktowy.

Zastosowanie rozwiązania minimaxowego gry strategicznej jako podstawy decyzji emitenta odzwierciedla postulat poszukiwania „decyzji najostrożniejszej”, tj. minimalizującej wartość maksymalnej straty; nie oznacza, że rynek jest w istocie graczem antagonistycznym.

Zasadniczym problemem zastosowania rozważanej gry w praktyce zarządzania długiem publicznym jest racjonalne określenie postaci macierzy gry $\mathbf{W}$. W niniejszej pracy zaproponowano dwa sposoby wyznaczenia tej macierzy, dostosowane do własności optymalnych portfeli długu.

2. Postać gry strategicznej

Przy założeniu, że dany jest zbiór strategii graczy (zbiór prognoz wariantowych) oraz sposób wyznaczenia optymalnego portfela długu (dla ustalonego wariantu prognozy), określenia wymaga jedynie postać elementów macierzy gry. Sformułowano dwie postacie tej macierzy, odpowiadające portfelom generującym identyczne lub nieidentyczne wpływy budżetowe.

Wyznaczenie optymalnego portfela instrumentów dłużnych $\mathbf{x}_s^*$ ($s=1, \dots, u$), odpowiadającego s -temu wariantowi prognozy, jest dokonywane na podstawie zadania minimalizującego koszty obsługi długu [4]. W niniejszej pracy przyjmuje się następującą postać tego zadania:
funkcja celu

$$\min_{\mathbf{x}} \left[\sum_{l=1}^K (M - d_s^{(l)}(x_l)) x_l \Psi_s^{(l)}(x_l) \right] = \min_{\mathbf{x}} \Phi_s(\mathbf{x}) \quad (1)$$

warunek określający poziom wpływów budżetowych (kapitału)

$$\sum_{l=1}^K x_l (M - d_s^{(l)}(x_l)) \geq a \quad (2)$$

gdzie:

x_l ($l=1, \dots, \kappa$) – sprzedaż l -tego instrumentu dłużnego (liczba nominalów)
– zmienna decyzyjna, κ ($\kappa \geq 2$) – liczba rodzajów emitowanych instrumentów, $\mathbf{x}$ – wektor zmiennych x_l ,

M – wartość nominalu instrumentu dłużnego,

$d_s^{(l)}(x_l)$ – średnie dyskonto przypadające na jeden nominal l -tego instrumentu dłużnego, odpowiadające sprzedaży przetargowej x_l nominalów oraz s -temu wariantowi stóp,

$\Psi_s^{(l)}(x_l)$ – składana stopa zwrotu [3, rozdz. 2, 3] l -tego instrumentu dłużnego, odpowiadająca: sprzedaży x_l nominalów oraz s -temu wariantowi prognozy stóp procentowych $r_{s,1}, \dots, r_{s,h}$, (h – horyzont prognozy), o postaci

$$\Psi_s^{(l)}(x_l) = \left((M * R_l \sum_{k=1}^{h-1} \prod_{j=k+1}^h (1 + r_{sj}) + M(1 + R_l)) / (M - d_s^{(l)}(x_l)) \right)^{1/h} - 1$$

przy czym: R_l – oprocentowanie kuponu l -tego instrumentu,

a – poziom wpływów budżetowych.

W przypadku skarbowych instrumentów dłużnych sprzedawanych na przetargach, funkcja celu (1) oraz funkcja występująca po lewej stronie nierówności (2) są funkcjami rosnącymi; ich własności omówiono w [4, rozdz. 4].

Oznaczając przez $\mathbf{x}_q^*$ ($q=1, \dots, u$) optymalne rozwiązanie zadania (1)-(2), dla q -tego wariantu prognozy stóp, symbol $\Phi_s(\mathbf{x}_s^*)$ wyraża średnioroczne koszty obsługi optymalnego portfela $\mathbf{x}_s^*$, w przypadku wystąpienia s -tego wariantu prognozy, a symbol $\Phi_s(\mathbf{x}_r^*)$ – średnioroczne koszty portfela $\mathbf{x}_r^*$, optymalnego dla r -tego wariantu prognozy, w przypadku wystąpienia s -tego wariantu prognozy stóp. Określenie wartości $\Phi_s(\mathbf{x}_r^*)$ opiera się na wariantowej prognozie stóp (zob. Klukowski 2003).

Przy założeniu, że każdy portfel $\mathbf{x}_q^*$ ($q=1, \dots, u$) generuje identyczne wpływy budżetowe, w przypadku wystąpienia dowolnego wariantu prognozy, tzn. są spełnione warunki

$$K_s(\mathbf{x}_r^*) = K_s(\mathbf{x}_s^*) \quad (r \neq s) \quad (3)$$

gdzie

$$K_s(\mathbf{x}_q^*) = \sum_{i=1}^{\kappa} x_{iq}^* (M - d_s^{(l)}(x_{iq}^*))$$

x_{iq}^* – l -ta składowa wektora $\mathbf{x}_q^*$,

zachodzi nierówność $\Phi_s(\mathbf{x}_r^*) \geq \Phi_s(\mathbf{x}_s^*)$, tj. koszty $\Phi_s(\mathbf{x}_r^*)$ są nie mniejsze niż optymalne koszty $\Phi_s(\mathbf{x}_s^*)$. Fakt ten implikuje sposób określenia elementów ω_{sr} macierzy gry $\mathbf{W}$, w postaci różnicy

$$\omega_{sr} = \Phi_s(\mathbf{x}_s^*) - \Phi_s(\mathbf{x}_r^*) \quad (s \neq r) \quad (4)$$

Każdy element ω_{sr} ($r \neq s$) wyraża stratę emitenta i wygraną drugiego gracza, tj. inwestorów, w wyniku sprzedaży przez emitenta portfela $\mathbf{x}_r^*$ oraz wystąpienia s -tego wariantu prognozy.

Założenie o równości wpływów budżetowych (równość (3)), ze sprzedaży poszczególnych portfeli $\mathbf{x}_1^*, \dots, \mathbf{x}_u^*$, może nie być spełnione. Jest tak dlatego, że dyskonto instrumentów dłużnych, określone przez wyniki przetargu, zależy od poziomu stóp procentowych; fakt ten wyjaśniono dokładniej [4, rozdz. 8]. Prawa strona równości (4) nie odzwierciedla wówczas dokładnie skutków błędów w wyborze wariantu prognozy, ponieważ koszty obsługi długu odpowiadające różnym wartościom wpływów budżetowych nie są porównywalne. Może to prowadzić, w szczególności, do dodatnich wartości ω_{sr} . W przypadku tym konieczna jest modyfikacja elementów macierzy gry. Można jej dokonać na różne sposoby; niżej proponuje się dokonać modyfikacji zadania służącego do konstrukcji optymalnego portfela instrumentów w taki sposób, aby zapewnić równość wpływów budżetowych, dla dowolnego wariantu prognozy. Polega ona na dodaniu do zadania (1)-(2) odpowiedniego warunku ograniczającego, w przypadku portfeli generujących nieidentyczne wpływy budżetowe.

Optymalny portfel $\tilde{\mathbf{x}}_r^*$ definiuje się obecnie w następujący sposób

$$\tilde{\mathbf{x}}_r^* = \begin{cases} \mathbf{x}_r^* & \text{w przypadku } K_s(\mathbf{x}_r^*) = K_s(\mathbf{x}_s^*); \\ \mathbf{x}_r^{(\min)} & \text{w przypadku } K_s(\mathbf{x}_r^*) < K_s(\mathbf{x}_s^*); \\ \mathbf{x}_r^{(\max)} & \text{w przypadku } K_s(\mathbf{x}_r^*) > K_s(\mathbf{x}_s^*), \end{cases} \quad (5)$$

gdzie:

- $\mathbf{x}_r^{(\min)}$ – rozwiązanie optymalne zadania (1)-(2) z dodanym warunkiem ograniczającym $\mathbf{x} \geq \mathbf{x}_r^*$, przy wykorzystaniu funkcji celu odpowiadającej s -temu wariantowi prognozy stóp procentowych,
- $\mathbf{x}_r^{(\max)}$ – rozwiązanie optymalne zadania (1)-(2) z dodanym warunkiem ograniczającym $\mathbf{x} \leq \mathbf{x}_r^*$, przy wykorzystaniu funkcji celu odpowiadającej s -temu wariantowi prognozy stóp procentowych.

Elementy $\tilde{\omega}_{sr}$ macierzy gry $\tilde{\mathbf{W}}$ odpowiadające zmodyfikowanemu portfelowi $\tilde{\mathbf{x}}_r^*$ definiuje się w sposób analogiczny do przypadku równości wpływów budżetowych, tj.

$$\tilde{\omega}_{sr} = \Phi_s(\mathbf{x}_s^*) - \Phi_s(\tilde{\mathbf{x}}_r^*) \quad (6)$$

Uzasadnienie sposobu konstrukcji portfela $\tilde{\mathbf{x}}_r^*$ jest następujące. Zaistnienie nierówności $K_s(\mathbf{x}_r^*) \neq K_s(\mathbf{x}_s^*)$ powoduje powiększenie zbioru warunków ograniczających zadania (1)-(2) o nierówności $\mathbf{x} \geq \mathbf{x}_r^*$ lub $\mathbf{x} \leq \mathbf{x}_r^*$ (dokładniej są to nierówności dla składowych wektora). Dodanie warunku $\mathbf{x} \geq \mathbf{x}_r^*$ implikuje optymalne zwiększenie składowych portfela $\mathbf{x}_r^*$ (zwiększenie sprzedaży) w taki sposób, aby zapewnić wpływy budżetowe na poziomie a , przy założeniu realizacji s -tego wariantu prognozy. Portfel otrzymany w ten sposób może być różny od portfela $\mathbf{x}_s^*$, co implikuje niedodatnią wartość wyrażenia (6). Analogicznie, warunek $\mathbf{x} \leq \mathbf{x}_r^*$ prowadzi do optymalnego zmniejszenia składowych wektora $\mathbf{x}_r^*$.

Obliczenie macierzy gry zdefiniowanej zależnością (6) jest bardziej pracochłonne niż macierzy określonej zależnością (4), ponieważ wymaga wielokrotnego rozwiązania zadania (1)-(2) z dodatkowym warunkiem ograniczającym.

Rozwiązanie minimaksowe określonej powyżej gry, z macierzą gry $\mathbf{W}$ lub $\tilde{\mathbf{W}}$, polega na wyznaczeniu optymalnej strategii emitenta. Jak już zaznaczono, ma ono postać rozkładu prawdopodobieństwa na zbiorze strategii (wariantów prognoz), wyznaczanego na podstawie zadania programowania liniowego. W przypadku, gdy macierz gry $\mathbf{W}'$ lub $\tilde{\mathbf{W}}'$ (symbole $\mathbf{W}'$, $\tilde{\mathbf{W}}'$ oznaczają macierze transponowane) zawiera punkt siodłowy rozkład ten jest jednopunktowy, co upraszcza rozwiązanie.

Punkt siodłowy γ macierzy $\mathbf{\Gamma} = [\gamma_{sr}]$ definiuje się zależnością

$$\gamma = \max_s (\min_r \gamma_{sr}) = \min_r (\max_s \gamma_{sr}) \quad (7)$$

Indeks s^* odpowiadający wartości γ , tzn. numer wiersza macierzy $\mathbf{\Gamma}$ zawierającego ten element, określa strategię stanowiącą optymalne rozwiązanie gry.

Optymalnym portfelem długu jest portfel $\mathbf{x}_{s^*}^*$ lub $\tilde{\mathbf{x}}_{s^*}^*$, odpowiadający punktowi siodłowemu macierzy gry – odpowiednio $\mathbf{W}'$ lub $\tilde{\mathbf{W}}'$.

Rozwiązanie takie minimalizuje stratę emitenta w przypadku wystąpienia najbardziej niekorzystnej dla niego ścieżki stóp procentowych – stanowi najostrożniejszą strategię w grze.

Rozwiązanie zrandomizowane ma postać funkcji prawdopodobieństwa ξ_r ($r=1, \dots, u$), $\sum_r \xi_r = 1$, określonej na zbiorze strategii emitenta (wariantów prognoz). Optymalna strategia zrandomizowana polega na wyborze jednego z wariantów prognozy oraz odpowiadającego mu optymalnego portfela instrumentów dłużnych w sposób losowy. Prawdopodobieństwa $\xi_1, \dots, \xi_u$ wyboru poszczególnych wariantów prognoz wyznacza się, w przypadku gry z macierzą $\mathbf{W}$, na podstawie rozwiązania zadania programowania liniowego o postaci

$$\min[1/v] = \min[p_1 + \dots + p_u] \quad (8)$$

$$\mathbf{W}'\mathbf{p} \geq \mathbf{1} \quad (9)$$

$$\mathbf{p} \geq \mathbf{0} \quad (10)$$

gdzie:

$p_r = \xi_r/v$, przy czym: $v = 1/(p_1, \dots, p_u)$; zależność ta służy do określenia prawdopodobieństw $\xi_r = v p_r$,

$\mathbf{p}$ – wektor (kolumnowy), którego elementami są wartości $p_1, \dots, p_u$,

$\mathbf{W}'$ – macierz $\mathbf{W}$ transponowana,

$\mathbf{1}$ – wektor jednostkowy (o składowych równych 1).

Losowy wybór optymalnego portfela $\mathbf{x}_s^*$ ($1 \leq s \leq u$), na podstawie funkcji prawdopodobieństwa $\xi_1, \dots, \xi_u$, może być dokonany przy wykorzystaniu generatora liczb losowych.

W przypadku gry z macierzą $\tilde{\mathbf{W}}$ należy dokonać odpowiedniej modyfikacji zadania (8)-(10), w szczególności zastąpić warunek (9) warunkiem $\tilde{\mathbf{W}}'\mathbf{p} \geq \mathbf{1}$.

3. Przykład zastosowania

Przedstawione wyżej koncepcje zastosowano do problemu minimalizacji ryzyka portfela obejmującego pięć instrumentów – obligacje stałoprocentowe: dwu-, pięcio-, i dziesięcioletnie, zmiennoprocentowe dziesięcioletnie oraz bony skarbowe 52-tygodniowe. Zbiór prognoz stóp procentowych (wyznaczono je w 2000 roku) obejmował 7 wariantów, podano je w tabeli 1.

Tabela 1

Wariantowa prognoza ścieżek rocznych stóp procentowych
na lata 2000-2010

Rok	Wariant						
	I	II	III	IV	V	VI	VII
2000	17,50	17,50	17,50	17,50	17,50	17,50	17,50
2001	13,00	16,69	17,62	16,00	15,37	13,00	15,37
2002	11,70	15,88	16,93	13,70	13,79	10,72	14,53
2003	10,70	15,08	16,17	11,70	12,37	9,73	13,40
2004	9,75	14,27	15,40	10,25	11,29	8,79	12,01
2005	9,25	13,46	14,51	9,75	10,85	8,30	10,85
2006	8,60	12,65	13,67	9,10	10,63	7,66	9,92
2007	8,15	11,85	12,77	8,40	10,49	7,22	9,50
2008	7,25	11,04	11,99	7,50	9,84	6,33	9,14
2009	7,00	10,23	11,04	7,25	9,10	6,09	9,10
2010	7,00	9,42	10,03	7,25	8,21	6,10	8,89

Na podstawie w/w prognoz oraz rzeczywistych wyników przetargów na skarbowe instrumenty dłużne wyznaczono siedem portfeli optymalnych x_q^* ($q=1, \dots, 7$), przy wykorzystaniu zadania (1)-(2). Wpływy budżetowe, odpowiadające tym portfelom, wykazują niewielkie różnice; dlatego też elementy macierzy gry określono, stosując pewne przybliżenia, według wzoru (4). Przykład zastosowania formuły (5) omówiono w [4, rozdz. 8].

Macierz strat $\mathbf{W}$, wyznaczona według wzoru (4), przyjęła postać

$$\mathbf{w} = 10^3 \begin{bmatrix} 0 & -12324 & -13853 & -8582 & -1106 & -1273 & -11285 \\ -19157 & 0 & -780 & -7855 & 0 & -33699 & 0 \\ -26277 & 0 & 0 & -10133 & 0 & -43612 & 0 \\ 0 & -1342 & -2262 & 0 & 0 & -342 & -235 \\ -3852 & -1485 & -2506 & -3601 & 0 & -10511 & -260 \\ -1839 & -20904 & -22515 & -14496 & -19745 & 0 & -19948 \\ -6403 & -1226 & -215 & -3591 & 0 & -14866 & 0 \end{bmatrix}$$

Rozwiązanie gry odpowiadającej powyższej macierzy $\mathbf{W}$ ma postać zrandomizowaną, ponieważ nie ma ona punktu siodłowego. Rozkład prawdopodobieństwa na zbiorze wariantów prognoz jest dwupunktowy, prawdopodobieństwa ξ_r są równe: $\xi_1 = \xi_2 = \xi_5 = \xi_6 = \xi_7 = 0$, $\xi_3 = 0,0422$, $\xi_4 = 0,9578$. Prawdopodobieństwo ξ_3 odpowiada wariantowi prognozy wykazującemu powolny spadek stóp, natomiast ξ_4 -wariantowi wykazującemu umiarkowany

spadek. W wyniku losowania (dokonanego przy wykorzystaniu generatora liczb losowych) otrzymano czwarty wariant prognozy, implikujący portfel optymalny o postaci

$$x_s^* = [4125000, 7326375, 1942720, 450000, 1089685]$$

Podsumowanie

Przedstawione wyżej koncepcje zastosowania dwuosobowej gry strategicznej do ograniczania ryzyka kosztów obsługi długu publicznego, wynikającego z wielowariantowości prognoz, stanowią składnik optymalizacyjnej metodologii zarządzania długiem publicznym, omówionej między innymi w [4-9]. Koncepcje te mogą być stosowane w warunkach rutynowej pracy podmiotu zarządzającego długiem, ze względu na prostotę teoretyczną oraz niewielkie wymagania obliczeniowe (wystarczający jest współczesny komputer PC oraz arkusz Excel); wzbogacają istniejące spektrum metod zarządzania ryzykiem. Stosowane narzędzi optymalizacyjnych w zarządzaniu długiem staje się koniecznością ze względu na obecną sytuację finansów publicznych w Polsce, tj. poziom deficytu budżetowego oraz dynamiczny wzrost zadłużenia publicznego i kosztów jego obsługi.

Literatura

1. Danish Government Borrowing and Debt (1999, 2000). Danmark Nationalbank, Copenhagen.
2. Greń J. (1972). Gry statystyczne i ich zastosowania. PWE, Warszawa.
3. Jajuga K., Jajuga T. (1997). Inwestycje, instrumenty finansowe, ryzyko finansowe, inżynieria finansowa. PWN, Warszawa.
4. Klukowski L. (2003). Optymalizacja decyzji w zarządzaniu instrumentami dłużnymi Skarbu Państwa. Wyższa Szkoła Informatyki Stosowanej i Zarządzania, Warszawa.
5. Klukowski L. (2005). Kompleksowa optymalizacja zarządzania zadłużeniem Skarbu Państwa. W: Zastosowania informatyki w nauce, technice i zarządzaniu. Red. J. Studziński, L. Drelichowski, O. Hryniewicz. PAN, Warszawa.
6. Klukowski L., Kuba E. (2001a). Optymalizacja zarządzania długiem Skarbu Państwa. Minimalizacja kosztów obsługi instrumentów dłużnych emitowanych na rynku krajowym. Materiały i Studia. NBP, Warszawa, z. 119.

7. Klukowski L., Kuba E. (2001b). Minimization of Public Debt Servicing Costs Based on Nonlinear Mathematical Programming Approach. *Control and Cybernetics*, 30.
8. Klukowski L., Kuba E. (2002a). Optymalizacja zarządzania długiem Skarbu Państwa w horyzoncie trzyletnim. W: *Badania Operacyjne i Systemowe wobec wyzwań XXI wieku. Modelowanie i optymalizacja, metody i zastosowania*. Red. J. Kacprzyk, J. Węglarz. Akademicka Oficyna Wydawnicza EXIT, Warszawa.
9. Klukowski L., Kuba E. (2002b). Stochastyczna optymalizacja strategii zarządzania skarbowymi instrumentami dłużnymi. *Materiały i Studia. NBP*, Warszawa, z. 152.
10. Peters H. (2008). *Game Theory. A Multi-Leveled Approach*. Springer-Verlag, Berlin, Heidelberg.

Honorata Sosnowska

Szkoła Główna Handlowa w Warszawie

Akademia Leona Koźmińskiego

REGUŁY FORMALNE RYZYKA PORZĄDKOWEGO W PRZYPADKU ZDARZEŃ POZYTYWNYCH I NEGATYWNYCH*

Wstęp

Zwykło się przyjmować, że z ryzykiem mamy do czynienia w sytuacji, gdy określamy prawdopodobieństwo mających się wydarzyć zdarzeń, zaś w sytuacji, gdy prawdopodobieństwa nie określamy, mamy do czynienia z niepewnością. Z pośrednią sytuacją mamy do czynienia, gdy wprowadzimy nie określamy liczbowej wartości ryzyka, ale stwierdzamy, że jakieś zdarzenie jest bardziej prawdopodobne od innego, czyli gdy porządkujemy zdarzenia pod względem ryzyka. Ta sytuacja znana jest jako ryzyko porządkowe (wymienne – prawdopodobieństwo porządkowe) i obejmuje bardzo wiele sytuacji, w których podejmujemy decyzje. Tak określane w bardzo rozmyty sposób prawdopodobieństwo jest często prawdopodobieństwem subiektywnym. To nasza opinia decyduje, czy coś uważamy za mniej lub bardziej prawdopodobne. Psychologowie społeczni uważają, że ludzie (decydenci) rzadko szacują dokładnie prawdopodobieństwo zdarzeń.

Celem niniejszej pracy jest zbadanie, czy prawa rachunku prawdopodobieństwa funkcjonują w przypadku prawdopodobieństwa porządkowego odnoszącego się do zdarzeń ewaluowanych pozytywnie, które w skrócie będziemy nazywać zdarzeniami pozytywnymi. W przypadku zdarzeń ewaluowanych negatywnie (w skrócie – negatywnych) prawa rachunku prawdopodobieństwa nie zawsze funkcjonują. To zostało zbadane w pracach Sosnowskiej [5; 6]. W niniejszej pracy porównane zostały również wyniki dotyczące zdarzeń pozytywnych i negatywnych. Znane są [2] niezgodności myślenia intuicyjnego i formalnego dotyczące koniunkcji (problem Lindy) i implikacji (Wasona

* Opracowanie to pod nazwą „Subiektywne prawdopodobieństwo w przypadku zdarzeń pozytywnie ewaluowanych” jest częścią badań statutowych KAE SGH nr 0008/08.

zadanie wyboru). W niniejszej pracy skupiliśmy się na alternatywie i porównaniu koniunkcji z alternatywą. Znaczna część badanych błędnie umieszcza porządkowo prawdopodobieństwo sumy zdarzeń określonej przez alternatywę zdań, ponadto w przypadku zdarzeń pozytywnych spora część błędnie określa porządek sumy i części wspólnej zdarzeń czyli myli koniunkcję i alternatywę.

Badania zostały przeprowadzone za pomocą eksperymentu kwestionariuszowego na 2 grupach studentów Szkoły Głównej Handlowej i 4 grupach studentów Akademii Leona Koźmińskiego (poprzednia nazwa – Wyższa Szkoła Przedsiębiorczości i Zarządzania im. Leona Koźmińskiego) różniących się stopniem przygotowania matematycznego.

Konstrukcja opracowania jest następująca: najpierw zostały przedstawione eksperymenty i zostały sformułowane hipotezy badawcze. Sformułowano główną hipotezę badawczą i hipotezy dodatkowe, następnie krótko przypomniano eksperyment dotyczący zdarzeń negatywnych, a dalej opisano eksperyment dotyczący zdarzeń pozytywnych. Porównano też wyniki dotyczące zgodności reguł formalnych z intuicyjnymi w przypadku obu typów zdarzeń. Połączono też aneksy gdzie zawarto użyte w eksperymentach formularze.

1. Eksperymenty

Opisane zostaną eksperymenty kwestionariuszowe dotyczące ryzyka porządkowego. Badani otrzymywali do wypełnienia kwestionariusz, w którym mieli uszeregować pod względem ryzyka podane zdarzenia. Następnie na tej skali, pomiędzy tymi zdarzeniami, mieli umieścić zdarzenia przeciwnie, sumę i część wspólną zdarzeń. Najbardziej prawdopodobne zdarzenia umieszczano najwyżej na skali i dalej kolejno coraz niżej aż do najmniej prawdopodobnych. Badania zostały przeprowadzone w dwu seriach na 6 grupach studentów. Dwie grupy były badane w przypadku zdarzeń negatywnych i cztery w przypadku zdarzeń pozytywnych. Wybór grup i ich liczebność uzależnione były tym, z jakimi studentami w tym czasie prowadzono zajęcia dydaktyczne. Grupom zostały nadane nazwy charakteryzujące kierunek studiów badanych studentów. W przypadku dwu grup z tego samego kierunku nadane im były numery. Wybór grup zapewniał ich rozłączność. Jako zdarzenia negatywne przyjęto zachorowania na wybrane choroby. Dobór chorób poprzedzono badaniem sondażowym mającym określić jakie choroby studenci w ogóle mogą brać pod uwagę. Podobnie postąpiono w przypadku zdarzeń pozytywnych. Zanim przystąpiono do konstrukcji kwestionariusza studenci w osobnym badaniu mieli podać zdarzenia niepewne o wydźwięku pozytywnym. Najczęściej wymieniane umieszczono w kwestionariuszu. W obu przypadkach badani nie znali dokładnego prawdopodobieństwa wystąpienia danego zdarzenia (nie zna się na ogół

zachorowalności na choroby ani np. szans dobrego ukończenia studiów czy wygranej w grze losowej) i nie mieli go określać. Mieli tylko odnieść się do używanej w swoim subiektywnym myśleniu relacji „mniejsze-większe”. Nie podano definicji chorób i zdarzeń pozytywnych (np. co to znaczy depresja czy dobry samochód). Badani odnosili się do prawdopodobieństwa subiektywnego. Kolejność przeprowadzenia eksperymentów była następująca. Najpierw zostały zbadane zdarzenia negatywne na grupie „Metody ilościowe II”, a potem na grupie „Psychologia II”. Potem zostały zbadane zdarzenia pozytywne mniej więcej w tym samym czasie na wszystkich czterech grupach („Socjologia”, „Zarządzanie”, „Psychologia I”, „Metody ilościowe I”). Pod względem przygotowania matematycznego grupy można uszeregować następująco. Najlepsze przygotowanie matematyczne miała grupa „Metody ilościowe I”, potem kolejno „Metody ilościowe II”, następnie „Zarządzanie”, „Psychologia II”, „Psychologia I” i na końcu „Socjologia”. Kolejność między obiema grupami z psychologii jest dość umowna. Studenci obu grup reprezentują zbliżony poziom. Umieściłam „Psychologię II” przed „Psychologią I” głównie z tego powodu, że studenci z grupy „Psychologia II” byli badani w trakcie wykładu z logiki i można przypuszczać, że dzięki temu lepiej pamiętali prawa logiczne niż studenci z grupy „Psychologia I”, którzy byli badani w 4 miesiące po egzaminie z logiki i mogli już sporo wiadomości zapomnieć.

1.1. Hipotezy badawcze

Celem eksperymentów było zbadanie, czy stosowane przez respondentów intuicyjne prawa prawdopodobieństwa porządkowego są zgodne z prawami formalnymi. Tego dotyczy główna hipoteza badawcza, hipoteza 1.

Hipoteza 1

Intuicyjnie stosowane prawa prawdopodobieństwa porządkowego w przypadku iloczynu i sumy zdarzeń różnią się od praw formalnych rachunku prawdopodobieństwa.

Następnie badano co powoduje, że te prawa się różnią. Sformułowano 3 hipotezy dodatkowe (hipotezy 2, 3, 4) dotyczące wpływu przygotowania matematycznego, tego jaki typ zdarzeń (pozytywne czy negatywne) badamy, tego o jakie prawo chodzi (dotyczące iloczynu czy sumy zdarzeń).

Hipoteza 2

Lepsze przygotowanie matematyczne zwiększa odsetek odpowiedzi zgodnych z prawami formalnymi rachunku prawdopodobieństwa.

Hipoteza 3

Odsetek odpowiedzi zgodnych z prawami formalnymi rachunku prawdopodobieństwa zależy od tego, czy badane zdarzenia są pozytywne czy negatywne.

Hipoteza 4

Odsetek odpowiedzi zgodnych z prawami formalnymi rachunku prawdopodobieństwa zależy od tego, którego działania na zdarzeniach one dotyczą, sumy czy iloczynu.

1.2. Zdarzenia negatywne

W eksperymentach użyto dwu wersji kwestionariusza (patrz kwestionariusze I i II w aneksie), dotyczącego ryzyka zachorowania na wybrane choroby. Konstrukcja obu wersji jest identyczna. Należy uporządkować pod względem ryzyka zachorowania podane choroby, a następnie na tej skali, pomiędzy zdarzeniami reprezentującymi choroby umieścić zdarzenia reprezentujące zdarzenia przeciwne, sumę i część wspólną zdarzeń. Kwestionariusze różnią się tylko jedną chorobą.

1. Grupa „Metody ilościowe” II

Eksperyment przeprowadzono na grupie 119 studentów pod koniec października 2006. Dokładny opis i analiza znajdują się w pracy [5], z której zaczerpnięto prezentowane w tym opracowaniu wyniki. Eksperymentowi zostali poddani studenci I roku studiów dziennych SGH, który uczęszczali na wykład z logiki matematycznej. Należy się spodziewać, że mieli dobre przygotowanie matematyczne, bo jednym z warunków przyjęcia na SGH jest zdanie matury z matematyki na rozszerzonym poziomie. Ponadto, przedmiot logika jest do wyboru w wersji opisowej i matematycznej, więc można się spodziewać, że na wykładzie z logiki matematycznej znajdą się ci, dla których matematyka nie stanowi problemu. W momencie eksperymentu studenci mieli za sobą wykłady z rachunku zdań i rachunku kwantyfikatorów. Użyty został kwestionariusz I (patrz aneks).

2. Grupa „Psychologia II”

Eksperyment dokładnie został opisany [6]. Eksperymentowi poddano 30 osób uczęszczających na wykład z logiki dla studentów I roku specjalności „Psychologia w zarządzaniu” studiów dziennych WSPiZ (obecna nazwa uczelni – Akademia Leona Koźmińskiego). Studenci mieli za sobą wykłady z matema-

tyki i psychologii, zaś na wykładzie z logiki omówiono wcześniej elementy rachunku zdań i kwantyfikatorów. Wykład z logiki jest obowiązkowy na tej specjalności. W rekrutacji do tej uczelni nie wymaga się specjalnego przygotowania matematycznego. Eksperyment przeprowadzony został wiosną 2007, mniej więcej pół roku po przeprowadzeniu eksperymentu w grupie „Metody ilościowe II”. W tym czasie wyniki pierwszego eksperymentu prezentowano i dyskutowano przed różnymi gremiami. W dyskusjach tych padła między innymi propozycja sprawdzenia, jak z badanym w eksperymencie zadaniem radzą sobie przyszli psychologowie. Zaproponowano również dodanie listy chorób jakiejś choroby występującej często, ale mogącej mieć istotne skutki dla zdrowia. Zdecydowano się zastąpić występującą w eksperymencie pierwszym biłałką złamaniem kończyny. W ten sposób powstał użyty w tym badaniu kwestionariusz II (patrz aneks).

1.3. Zdarzenia pozytywne

Na konferencji „Psychologia Ekonomiczna” we Wrocławiu wiosną 2008 przedstawiono wyniki dotyczące formalnych i intuicyjnych reguł rachunku prawdopodobieństwa w przypadku zdarzeń negatywnych. Padło wtedy pytanie czy rezultaty będą takie same w przypadku zdarzeń pozytywnych. W odpowiedzi na to pytanie został skonstruowany kwestionariusz dotyczący zdarzeń pozytywnych (kwestionariusz III w aneksie) i przeprowadzona seria eksperymentów na 4 grupach studentów w kwietniu i maju 2008. Większa liczba grup niż w przypadku zdarzeń negatywnych wynikała po pierwsze z ich mniejszej liczebności, a ponadto z chęci sprawdzenia jak radzą sobie z zadaniami eksperymentalnymi badani, którzy w ogóle nie mieli matematyki, co spowodowało dołączenie grupy „Socjologia”. Ponieważ wykład dla socjologii i zarządzania był wspólny, więc eksperymentowi poddano także studentów zarządzania.

1. Grupa „Socjologia”

Grupa składała się z ośmiu studentów studiów dziennych I roku socjologii WSPiZ (obecna nazwa uczelni – Akademia Leona Koźmińskiego) uczęszczających na wykład z logiki. Wykład był wspólny dla studentów socjologii i zarządzania więc eksperyment przeprowadzono wspólnie. Wyniki jednak zostały podzielone na dwie grupy, w zależności od kierunku studiów po to, aby uwzględnić różne przygotowanie matematyczne badanych. Studenci socjologii nie mają w programie studiów żadnych zajęć z matematyki, podczas gdy studenci zarządzania mają. Wykład z logiki był dla studentów socjologii obowiązkowy.

2. Grupa „Zarządzanie”

Grupa składała się z 9 studentów pierwszego roku studiów dziennych na kierunku zarządzanie WSPiZ (obecna nazwa uczelni – Akademia Leona Koźmińskiego) na wspólnym dla zarządzania i socjologii wykładzie z logiki na I roku studiów. Wykład dla studentów zarządzania był do wyboru z pary logika – historia gospodarcza. Biorąc pod uwagę fakt, że wybrali logikę można przypuszczać, że na wykład trafili studenci stosunkowo dobrze przygotowani matematycznie. Studenci zarządzania uczęszczali ponadto na obowiązkowy wykład z matematyki.

3. Grupa „Psychologia I”

Grupa składała się z 38 studentów studiów dziennych pierwszego roku specjalności psychologia w zarządzaniu WSPiZ (obecna nazwa uczelni – Akademia Leona Koźmińskiego). Studenci w poprzednim semestrze mieli obowiązkowy wykład z logiki. Eksperyment został przeprowadzony na wykładzie ze statystyki. Studenci mieli też obowiązkowe zajęcia z matematyki.

4. Grupa „Metody ilościowe I”

Grupa składała się z 27 studentów studiów dziennych kierunku metody ilościowe i systemy informacyjne w ekonomii SGH. Eksperyment przeprowadzono na wykładzie z ekonomii matematycznej przeznaczonym dla studentów III roku i wyższych lat studiów. Wykład ten jest wykładem do wyboru. Studenci uczęszczali poprzednio na liczne przedmioty matematyczne. Ze wszystkich badanych grup jedynie ta miała bardzo dobre przygotowanie matematyczne.

2. Reguły formalne a reguły intuicyjne

Wyniki badania zdarzeń negatywnych [5; 6] wskazują, że badani mieli kłopot z poprawnym umieszczeniem na skali (wyżej lub na tej samej wysokości co poszczególne zdarzenia) sumy zdarzeń. Tym niemniej, znakomita większość umieszczała sumę zdarzeń jako bardziej prawdopodobną (czyli wyżej na skali) niż część wspólna zdarzeń. W badaniu dotyczącym zdarzeń pozytywnych postanowiono sprawdzić hipotezę, że badani poprawnie umieszczają na skali prawdopodobieństwo sumy i części wspólnej zdarzeń. Wyniki przedstawia tabela 1. W kolumnie trzeciej, oznaczonej „ $\wedge \leq \vee$ ”, umieszczono dane dotyczące wyników poprawnych, zaś w kolumnie czwartej, oznaczonej „ $\vee < \wedge$ ” – niepoprawnych. W przypadku grup o najsłabszym przygotowaniu matematycz-

nym („Socjologia”, „Psychologia I”) hipoteza o prawidłowym umieszczaniu na skali prawdopodobieństwa części wspólnej i sumy zdarzeń się nie potwierdziła. Można wręcz twierdzić, że badani mają kłopoty z odróżnieniem koniunkcji i alternatywy zdań.

Tabela 1

Zależności pomiędzy prawdopodobieństwem części wspólnej i sumy zdarzeń pozytywnych
(w nawiasach okrągłych podano wyniki w procentach liczone dla badania opisanego
w danym wierszu)

Grupa	Liczba badanych	$\wedge \leq \vee$	$\vee < \wedge$
1.Socjologia	8	4 (50%)	4 (50%)
2. Zarządzanie	9	6 (66%)	3 (33%)
3.Psychologia I	38	12 (31%)	26 (68%)
4. Metody ilościowe I	27	23 (85%)	4 (15%)

Następnie sprawdzono jak w obu badaniach respondenci sytuują odpowiednio część wspólną i sumę zdarzeń w stosunku do tych zdarzeń. Wyniki dotyczące zdarzeń pozytywnych przedstawia tabela 2. Sprawdzano, czy badani umieszczają część wspólną i sumę zdarzeń jako mniej prawdopodobne od obu zdarzeń (kolumna trzecia, oznaczona $\wedge \leq \dots, \dots; \vee \leq \dots, \dots$), czy lokują pomiędzy zdarzeniami (kolumna czwarta, oznaczona $\dots < \wedge < \dots; \dots < \vee < \dots$), czy też uważają za bardziej prawdopodobne od obu zdarzeń (kolumna piąta, oznaczona $\dots, \dots \leq \wedge; \dots, \dots \leq \vee$). Przypadki, gdy traktowane prawdopodobieństwo sumy lub części wspólnej zdarzeń jako równe prawdopodobieństwu któregoś ze zdarzeń były bardzo nieliczne. Jak widać poprawnie (kolumna trzecia, oznaczona $\wedge \leq \dots, \dots$, w przypadku części wspólnej i kolumna piąta, oznaczona $\dots, \dots \leq \vee$, w przypadku sumy zdarzeń) umieszczało zdarzenia mniej niż 50% badanych z grup o słabszym przygotowaniu matematycznym („Socjologia”, „Zarządzanie”, „Psychologia I”). Warto zauważyć, że również w grupie o bardzo dobrym przygotowaniu matematycznym „Metody ilościowe I” spora część badanych dała błędna odpowiedź. Błędne odpowiedzi w przypadku części wspólnej zdarzeń można wyjaśnić występowaniem zauważonego przez Kahnemana i Tverskiego problemu Lindy [2; 4]. Jeśli zaś chodzi o błędne umieszczanie sumy zdarzeń, to badanie potwierdza wyniki uzyskane dla zdarzeń negatywnych, że pojęcie sumy zdarzeń określanej przez alternatywę zdań jest mało intuicyjne.

Tabela 2

Zależności pomiędzy prawdopodobieństwem części wspólnej sumy zdarzeń pozytywnych i prawdopodobieństwem tych zdarzeń (w nawiasach okrągłych podano wyniki w procentach liczone dla badania opisanego w danym wierszu)

Grupa – pozytywne	Liczba badanych	$\wedge \leq \dots, \dots$	$\dots < \wedge < \dots$	$\dots, \dots \leq \wedge$
1. Socjologia	8	3 (37%)	4 (50%)	1 (12%)
2. Zarządzanie	9	1 (11%)	6 (66%)	2 (22%)
3. Psychologia I	38	2 (5%)	20 (53%)	16 (42%)
4. Metody ilościowe I	27	13 (48%)	10 (37%)	4 (15%)
Grupa – pozytywne	liczba badanych	$\vee \leq \dots, \dots$	$\dots < \vee < \dots$	$\dots, \dots \leq \vee$
1. Socjologia	8	3 (37%)	5 (62%)	0
2. Zarządzanie	9	4 (44%)	3 (33%)	2 (22%)
3. Psychologia I	38	12 (31%)	18 (47%)	8 (21%)
4. Metody ilościowe I	27	2 (7%)	10 (37%)	15 (55%)

Wyniki analogicznej analizy dla zdarzeń negatywnych zawiera tabela 3. Zauważmy, że w przeciwieństwie do przypadku zdarzeń pozytywnych znakomita większość badanych właściwie sytuuje część wspólną zdarzeń. Natomiast wyniki dotyczące sumy zdarzeń pokazują, że większość sytuuje ją błędnie, przy czym wyniki są zbliżone do przypadku zdarzeń pozytywnych. Gorsze wyniki w przypadku zdarzeń pozytywnych niż negatywnych można próbować wyjaśnić stosunkiem do ryzyka. Ryzyko zdarzenia negatywnego odbieramy jako ryzyko zagrożenia i pewno dokładniej analizujemy. Jeśli zaś chodzi o ryzyko zdarzeń pozytywnych, to przy błędnym szacunku najwyżej tego dobrego zdarzenia nie będzie, ale nic złego się nie zdarzy, stąd można mniej dokładnie analizować ryzyka zdarzeń pozytywne.

Tabela 3

Zależności pomiędzy prawdopodobieństwem części wspólnej i sumy zdarzeń negatywnych oraz prawdopodobieństwem tych zdarzeń (w nawiasach okrągłych podano wyniki w procentach liczone dla badania opisanego w danym wierszu)

Grupa – negatywne	Liczba badanych	$\wedge \leq \dots, \dots$	$\dots < \wedge < \dots$	$\dots, \dots \leq \wedge$
1. Metody ilościowe II	119	100 (84%)	15 (13%)	4 (3%)
2. Psychologia II	28	25 (89%)	1 (4%)	2 (7%)
Grupa – Negatywne	liczba badanych	$\vee \leq \dots, \dots$	$\dots < \vee < \dots$	$\dots, \dots \leq \vee$
1. Metody ilościowe II	119	41 (35%)	35 (30%)	41 (35%)
2. Psychologia II	26	15 (57%)	7 (27%)	4 (15%)

W tabeli 4 przedstawiono zależności dotyczące zdarzenia i zdarzenia przeciwnego. Pierwsze cztery grupy to badanie zdarzeń pozytywnych. Większość badanych przewiduje, że jest bardziej prawdopodobne, iż będą dobrze zarabiać niż że nie będą dobrze zarabiać (trzecia kolumna), co świadczy o ich optymizmie. Wyniki z czwartej kolumny można interpretować jako racjonalność badanych, znakomita mniejszość ocenia wygraną 100.000 zł w totolotka jako bardziej prawdopodobną niż brak takiej wygranej. Trzy ostatnie kolumny dotyczą zjawisk negatywnych. W obu grupach badani mieli zająć prawdopodobieństwem zachorowania na depresję. Znakomita większość uznała za bardziej prawdopodobne, że zachorują na depresję niż że nie zachorują. Wynika to pewno z szerokiego traktowania znaczenia słowa depresja, jak i z sygnalizowanego ostatnio przez media stanu ducha młodego pokolenia. Kolejne dwie kolumny mogą służyć za przykład zagadnienia zwanego przeszacowywaniem zjawisk rzadkich [1]. Przechodząc do interpretacji liczbowej 19% badanych uznało, że z prawdopodobieństwem co najmniej 50% zachoruje na białaczkę, zaś 85%, że złamie kończynę. Jest to znacznie więcej niż zachorowalność na te choroby.

Tabela 4

Prawdopodobieństwo wybranych zdarzeń i zdarzeń przeciwnych do nich

Grupa	Liczba badanych	$\neg Z \leq Z$ pozytywne	$\neg W \leq W$ pozytywne	$\neg D \leq D$ negatywne	$\neg B \leq B$ negatywne	$\neg ZK \leq ZK$ negatywne
1. Socjologia	8	4 (50%)	2 (25%)	*	*	*
2. Zarządzanie	9	7 (77%)	2 (22%)	*	*	*
3. Psychologia I	38	33 (87%)	5 (13%)	*	*	*
4. Metody ilościowe I	27	26 (96%)	1 (3%)	*	*	*
5. Metody ilościowe II	119	*	*	75 (64%)	23 (19%)	*
6. Psychologia II	30	*	*	22 (73%)	*	25 (85%)

Oznaczenia: Z – będę dobrze zarabiać, W – wygram 100.000 w totolotka, D – zachoruję na depresję, B – zachoruję na białaczkę, ZK – złamię kończynę. $\neg$ – zaprzeczenie.

Podsumowanie

Wyniki przedstawione w tabelach 1, 2, 3 pokazują, że tylko część badanych poprawnie określa prawdopodobieństwo sumy i iloczynu zdarzeń wzajemnie względem siebie i względem zdarzeń, z których zostały skonstruowane (kolumna 3 w tabeli 1 oraz kolumna 3 w przypadku iloczynu i kolumna 5 w przypadku sumy w tabelach 2 i 3). W niektórych grupach badanych jest to poniżej 50%. W ten sposób główna hipoteza badawcza – hipoteza 1, o braku zgodności intuicyjnych praw prawdopodobieństwa porządkowego z prawami

formalnymi została potwierdzona. Hipotezy dodatkowe próbują wyjaśnić przyczyny tego braku. W grupach o lepszym przygotowaniu matematycznym (Metody Ilościowe I i II) procent poprawnych odpowiedzi jest na ogół wyższy niż w pozostałych grupach. Dokładna analiza statystyczna wykonana przez panią Monikę Kozińską daje umiarkowaną dodatnią Korelację Spearmana na poziomie istotności 0,05 między poprawnością odpowiedzi i przygotowaniem matematycznym dla zdarzeń pozytywnych. W przypadku zdarzeń negatywnych jest słaba dodatnia korelacja na poziomie 0,1 dla iloczybu i brak podstaw dla odrzucenia hipotezy o braku korelacji dla iloczynu. Potwierdza hipotezę 2 o wpływie przygotowania matematycznego na poprawność odpowiedzi w przypadku zdarzeń pozytywnych i częściowo acz słabiej w przypadku negatywnych. Porównanie odsetka poprawnych odpowiedzi w przypadku zdarzeń pozytywnych i negatywnych pokazuje, że różnią się one istotnie tylko w przypadku iloczynu zdarzeń, gdzie więcej poprawnych uzyskano dla zdarzeń negatywnych. Można to tłumaczyć tak, że w przypadku zdarzeń negatywnych nie działa optymizm i badani staranniej analizują swoje odpowiedzi. Hipoteza 3 została potwierdzona tylko częściowo.

Istotnie większy odsetek poprawnych odpowiedzi w przypadku iloczynu niż w przypadku sumy daje się zauważyć tylko dla zdarzeń negatywnych. Tak więc hipoteza 4 została potwierdzona tylko częściowo.

Konkludując, intuicyjne subiektywne prawa prawdopodobieństwa porządkowego nie pokrywają się z prawami rachunku prawdopodobieństwa. Im lepsze przygotowanie matematyczne, tym większa zbieżność praw intuicyjnych i formalnych, zwłaszcza w przypadku zdarzeń pozytywnych.

ANEKS

KWESTIONARIUSZ I

Bierze Pani/Pan udział w badaniu dotyczącym ryzyka porządkowego.
Badanie składa się z dwu części.

CZĘŚĆ I

Następujące choroby trzeba umieścić w wierszach oznaczonych numerami od 1 do 6 w zależności od ryzyka zachorowania na tę chorobę, zaczynając od tej, na którą ryzyko zachorowania jest wg. Pani/Pana największe, kolejno te o coraz mniejszym ryzyku, a kończąc na tej, na którą ryzyko zachorowania jest najmniejsze.

CHOROBY

AIDS, białaczka, depresja, ptasia grypa, rak, schizofrenia

.....
1.....
.....
2.....
.....
3.....
.....
4.....
.....
5.....
.....
6.....
.....

CZĘŚĆ II

W wierszach, które nie są oznaczone numerami, należy w zależności od ocenianego przez Panią/Pana ryzyka wystąpienia takiej sytuacji umieścić następujące zdarzenia:

- a) nie zachoruje na białaczkę
- b) nie zachoruje na depresję
- c) zachoruje na białaczkę i depresję
- d) zachoruje na AIDS lub raka

KWESTIONARIUSZ II

Bierze Pani/Pan udział w badaniu dotyczącym ryzyka porządkowego.

Badanie składa się z dwu części.

CZEŚĆ I

Następujące choroby trzeba umieścić w wierszach oznaczonych numerami od 1 do 6 w zależności od ryzyka zachorowania na tę chorobę, zaczynając od tej, na którą ryzyko zachorowania jest wg. Pani/Pana największe, kolejno te o coraz mniejszym ryzyku, a kończąc na tej, na którą ryzyko zachorowania jest najmniejsze.

CHOROBY

AIDS, złamanie kończyny, depresja, ptasia grypa, rak, schizofrenia

.....
1.....
.....
2.....
.....
3.....
.....
4.....
.....
5.....
.....
6.....
.....

CZEŚĆ II

W wierszach, które nie są oznaczone numerami, należy w zależności od ocenianego przez Panią/Pana ryzyka wystąpienia takiej sytuacji umieścić następujące zdarzenia:

- e) nie złamię kończyny
- f) nie zachoruję na depresję
- g) zachoruję na AIDS i raka
- h) zachoruję na AIDS lub raka

KWESTIONARIUSZ III

RYZIKO PORZĄDKOWE – ZDARZENIA POZYTYWNE

Kierunek studiów....., uczelnia

Biorą Państwo udział w eksperymencie dotyczącym ryzyka porządkowego.

Eksperyment składa się z dwu części.

CZĘŚĆ I.

Niżej zostały podane w porządku alfabetycznym zdarzenia pozytywne, ale niepewne. Należy je wpisać w wierszach oznaczonych liczbami 1-6 tak, aby uporządkować od najbardziej prawdopodobnego (umieszczonego w wierszu 1) do najmniej prawdopodobnego (umieszczonego w wierszu 6). Te zdarzenia to

Mąż/Żona – znajdę dobrego współmałżonka

Podróże – będę dużo podróżować

Samochód – będę mieć dobry samochód

Studia –dobrze ukończę studia

Wygrana – wygram 100 000 zł w totolotka

Zarobki – będę dobrze zarabiać

.....

1.....

.....

2.....

.....

3.....

.....

4.....

.....

5.....

.....

6.....

.....

CZĘŚĆ II

W wierszach nieoznaczonych liczbami umieść następujące zdarzenia

(a) *Nie wygram 100 000 zł w totolotka*

(b) *nie będę dobrze zarabiać*

(c) *będę mieć dobry samochód i dobrze ukończę studia*

(d) *będę mieć dobry samochód lub dobrze ukończę studia*

Literatura

1. Bernstein, Peter L. (1997). *Przeciw bogom. Niezwykłe dzieje ryzyka*. WIG-Press, Warszawa.
2. Chiswell I., Hodges W. (2007). *Mathematical Logic*. Oxford University Press, Oxford.
3. Grupowe podejmowanie decyzji. Elementy teorii. Przykłady zastosowań. (1999). Red. H. Sosnowska. Wydawnictwo Naukowe SCHOLAR, Warszawa.
4. Nęcka E., Orzechowski J., Szymura B. (2006). *Psychologia poznawcza*. PWN, Warszawa.
5. Sosnowska H. (2008a). Ryzyko porządkowe a rachunek zdarzeń. W: *Modelowanie preferencji a ryzyko '08*. Red. T. Trzaskalik. AE, Katowice, 91-100.
6. Sosnowska H. (2008b). Prawdopodobieństwo subiektywne a reguły rachunku prawdopodobieństwa. W: *Metody ilościowe w ekonomii*. Wyższa Szkoła Bankowa, Poznań, 2-22.

Maciej Wolny

Politechnika Śląska w Gliwicach

QUASI-KLASYCZNE METODY WSPOMAGANIA RACJONALNEGO PODZIAŁU WYGRANEJ W KOOPERACYJNYCH GRACH Z KOALICJAMI

Wprowadzenie

Jednym z podstawowych problemów związanych z wieloosobowymi grami z koalicjami jest zagadnienie ustalenia podziału wygranej między członków koalicji. Ustalony podział powinien być akceptowany przez wszystkich graczy tworzących koalicję.

Inspiracją do podjęcia tematu jest zagadnienie wspomagania decyden-
ta (rozumianego również jako kolektyw) w sytuacjach, gdy przyjęty ad hoc sposób podziału wygranej, uznany za sprawiedliwy, okazuje się nieracjonalny, a tym samym nie do zaakceptowania przez gremium. Głównym celem pracy jest przedstawienie wybranych, najczęściej rozpatrywanych metod podziału wygranej (a w przypadku quasi-klasycznych, metod wspomagania podziału) między graczy w grach z koalicjami, przedstawienie przyczyn nieracjonalności niektórych klasycznych podziałów oraz uzasadnienia stosowania metod opartych na logice programowania wielokryterialnego.

Podział wygranej powinien być racjonalny, w tym sensie, że spełniania warunki racjonalności indywidualnej, zbiorowej i koalicyjnej – znajduje się w rdzeniu gry [17].

Do klasycznych koncepcji ustalenia podziału wygranej, w sensie niniejszej pracy, zaliczane będą wszystkie metody normatywnie wskazujące podział „z góry”, przyjmujące pewną logikę, lub aksjomaty, które powinien spełniać podział. Koncepcjami klasycznymi analizowanymi w pracy są: wartość Shapleya [14], wartość Banzhafa [2], Nucleolus [16], punkt Gatela'yego [6].

Natomiast za metody quasi-klasyczne uważane będą metody ustalenia podziału wygranej bazujące na logice programowania wielokryterialnego, umożliwiające przyjęcie parametrów, od których może zależeć rozwiązanie. Przedrostek „quasi-” jest zasadny o tyle, że metody programowania umożliwiają znalezienie podziału, który spełnia z góry narzucone na podział założenia

(dotyczące w szczególności spełnienia odpowiednio zdefiniowanych warunków racjonalności). Można jednak zauważyć, że istnieją metody zaliczane do klasycznych, które mogą jednakowo realizować rozumowanie zdefiniowane jako quasi-klasyczne (koncepcja nucleolusa i punkt Gatelly'ego). Koncepcje zdefiniowane jako quasi-klasyczne (w odniesieniu do rozwiązania zagadnienia ustalenia podziału wygranej w grach z koalicjami), które przedstawia opracowanie to metody: programowanie leksykograficzne [1], skalaryzacji przez zastosowanie sumy ważonej [8], punktu referencyjnego [19] i programowania celowego. Na rozważane zagadnienie można również spojrzeć w szerszym znaczeniu zagadnienia programowania leksykograficznego [8].

W pracy nie skupiono szczególnej uwagi na innych koncepcjach, które prezentują różne ujęcia przedstawionych koncepcji [7; 11; 13; 18], rozwiązania Raiffa (Kalai-Smorodinskiego) [9, 12] oraz interaktywnej procedury wielokryterialnego wspomaganie podziału [4]. Jednak prace te również były inspiracją po podjęcia rozważanego zagadnienia.

1. Podstawowe definicje i założenia

Kooperacyjną grę z koalicjami w postaci funkcji charakterystycznej rozumie się jako uporządkowaną parę $\Gamma = (N, v)$, gdzie: $N = \{1, 2, \dots, n\}$ jest zbiorem numerów graczy, v jest funkcją charakterystyczną gry, która przyporządkowuje każdemu podzbiorowi $S \subseteq N$ łączną maksymalną nieujemną wypłatę $v(S)$. Przy tym $v(\emptyset) = 0$ oraz v jest superaddytywna, czyli

$$v(S) + v(T) \leq v(S \cup T), \quad S, T \subset N, \quad S \cap T = \emptyset. \quad (1)$$

Powyższe założenie (superaddytywności funkcji charakterystycznej) skutkuje tym, że wspólne działanie graczy jest zawsze (dla dowolnego zbioru graczy) nie gorsze niż działanie samodzielne lub w mniejszych koalicjach.

Rozważania dotyczą gier z wypłatami ubocznymi, czyli gier w których gracze mogą dzielić się wypłatami. Ponadto, rozważane będą gry istotne, czyli wygrana wielkiej koalicji (składającej z wszystkich graczy) musi być większa niż suma wygranych graczy w sytuacji, gdy działają pojedynczo (w koalicjach jednosobowych)

$$v(N) > \sum_{i=1}^n v(\{i\}) \quad (2)$$

Podziałem wygranej między graczy w grze niech będzie wektor

$$x = (x_1, \dots, x_n) \in X \subseteq \mathbb{R}_+^n \quad (3)$$

który określa wypłaty graczy w danej (przez zawiązanie koalicji i negocjacje) sytuacji. X jest zbiorem wszystkich możliwych podziałów wypłat w grze.

Podział spełnia postulat racjonalności zbiorowej, gdy spełniony jest warunek

$$\sum_{i=1}^n x_i = v(N) \quad (4)$$

Dodatkowo, gdy spełniony jest warunek racjonalności indywidualnej

$$x_i \geq v(\{i\}), \quad i = 1, 2, \dots, n \quad (5)$$

podział nazywany jest imputacją. Natomiast, warunek racjonalności koalicyjnej ma następującą postać

$$\sum_{i \in S} x_i \geq v(S), \quad S \subseteq N \quad (6)$$

Zbiór podziałów spełniający jednocześnie wszystkie trzy warunki racjonalności nazywany jest rdzeniem gry. Każdy podział nienależący do rdzenia będzie uznawany za nieracjonalny, a sytuacja, gdy rdzeń będzie zbiorem pustym, będzie dotyczyła gier, w których nie jest możliwy racjonalny podział. W pracy nie podjęto zagadnienia dotyczącego takich sytuacji, tym samym przyjmuje się, że istnieje niepusty rdzeń gry.

Warunki na niepusty rdzeń zostały sformułowane niezależnie przez Bondareva [3] i Shapleya [15] i związane są z liniowymi ograniczeniami dotyczącymi rdzenia gry. W grze istnieje niepusty rdzeń, jeśli równocześnie spełnione są następujące warunki

$$v(N) \geq \sum_{S \subseteq N \setminus \{\emptyset\}} \lambda_S v(S), \quad \lambda_S \geq 0, \quad S \subseteq N \setminus \{\emptyset\} \quad (7)$$

$$\sum_{S \subseteq N \setminus \{\emptyset\}} \lambda_S a_S = a_N \quad (8)$$

gdzie λ_S jest nieujemną liczbą rzeczywistą, natomiast a_S jest wektorem charakterystycznym koalicji S takim, że i -ta składowa jest równa 1 jeśli i -ty gracz należy do koalicji S oraz 0 jeśli nie należy.

Problem, będący punktem wyjścia do rozważań, polega na ustaleniu wygranych graczy tworzących koalicję, która się zawiąże. W pracy skupiono się na sytuacji, w której zawiązuje się jedna wielka koalicja składająca się z wszystkich graczy.

2. Metody (wspomagania) podziału

Aksjomatyczną koncepcją sprawiedliwego podziału wygranej między graczy jest wartość Shapleya [14]

$$\varphi_i = \frac{1}{n!} \sum_{\substack{S \subseteq N \\ i \in S}} (s-1)!(n-s)! [v(S) - v(S - \{i\})] \quad (9)$$

którą można interpretować jako średnią wartość jaką wnosi do wielkiej koalicji gracz nr i , w przypadku, gdy kolejność tworzenia tej koalicji jest równoprawdopodobna, s jest liczebnością koalicji S . Wartość Shapleya $\varphi = (\varphi_1, \dots, \varphi_n)$ jest zawsze imputacją.

Kolejnym sposobem określenia podziału jest wykorzystanie wartości Banzhafa [2], która dla i -tego gracza wynosi

$$\beta_i = \frac{1}{2^{n-1}} \sum_{S \subseteq N \setminus \{i\}} [v(S \cup \{i\}) - v(S)] \quad (10)$$

Wartość Banzhafa $\beta = (\beta_1, \dots, \beta_n)$ jest wykorzystywana jako miara znaczenia (siły) graczy w koalicji, stąd koncepcja podziału wygranej według, tak rozumianego, znaczenia graczy w koalicji, można ją interpretować jako wartość jaką wnosi do koalicji gracz, będąc w danej koalicji (w odróżnieniu od wartości Shapleya, która mówi o wartości przy wejściu do koalicji). Wartość ta jednak zwykle nie jest imputacją i nie należy do rdzenia gry. Dlatego więc, aby spełniała warunek racjonalności zbiorowej, często poddaje się ją normalizacji

$$\beta_i^* = \frac{\beta_i \cdot v(N)}{\sum_{i \in N} \beta_i} \quad (11)$$

Kolejną metodą, którą można zastosować do wyznaczenia imputacji jest rozwiązanie Raiffa $\rho = (\rho_1, \dots, \rho_n)$ [12], które można wyznaczyć dla i -tego gracza według następującej formuły [19 s. 85]:

$$\rho_i = x_i^l + \frac{v(N) - \sum_{i=1}^n x_i^l}{\sum_{i=1}^n (x_i^u - x_i^l)} (x_i^u - x_i^l) \quad (12)$$

gdzie $x_i^l = v(\{i\})$ jest dolnym ograniczeniem rdzenia gry dla i -tego gracza, natomiast $x_i^u = v(N) - v(N - i)$ górnym ograniczeniem. Koncepcja rozwiązania Raiffa wywodzi się z analizy problemu przetargu i ma charakter aksjomatyczny [9]. Wartość ta ma szczególne znaczenie jako jeden z polecanych punktów referencyjny w koncepcji Wierzbickiego przedstawionej w dalszej części pracy.

Kolejną metodą wyodrębnienia jednego podziału, jest koncepcja nukleolusa zaproponowana przez Schmeidlera [16]. Koncepcja ta, w odniesieniu do rozpatrywanej klasy gier, realizuje postulat maksymalizacji największego zadowolenia z koalicji [17; s. 261]. Definiuje się przekroczenie jako miarę niezadowolenia z koalicji S

$$e_S(x) = v(S) - \sum_{i \in S} x_i \quad (13)$$

która w przypadku niepustego rdzenia jest zawsze niedodatnia, wobec tego maksymalizuje się następujące wyrażenie

$$\min_S |e_S(x)| \rightarrow \max \quad (14)$$

W sytuacji, gdy wyrażenie (14) jest maksymalne dla więcej niż jednego podziału, należy maksymalizować drugie w kolejności minimum, następnie trzecie, i tak dalej, aż otrzymany zostanie jednoelementowy zbiór podziałów. Nukleolus zawsze należy do rdzenia gry oraz jest pojedynczą imputacją.

Następnym sposobem wyznaczenia pojedynczego podziału jest koncepcja punktu Gately'ego [5], która z kolei realizuje postulat minimalizacji maksymalnej skłonności do zerwania wielkiej koalicji. Miara skłonności do zerwania koalicji jest zdefiniowana przez następujące wyrażenie

$$d_i(x) = \frac{\sum_{j \neq i} x_j - v(N - i)}{x_i - v(i)} \quad (15)$$

Poszukuje się więc podziału spełniającego kryterium

$$\max_{i \in N} d_i(x) \rightarrow \min \quad (16)$$

Przedstawione powyżej metody umożliwiają wybór pojedynczej imputacji.

Punktem wyjścia do rozważań podjętego zagadnienia przy zastosowaniu programowania wielokryterialnego jest stwierdzenie, że każdy z racjonalnych graczy chce otrzymać możliwie najwyższą wypłatę – każda ze składowych wektora x jest maksymalizowana [1]

$$x \rightarrow \max \quad (17)$$

przy jednoczesnym spełnieniu postulatów racjonalności (4)-(6).

Rozwiązanie zagadnienia dotyczącego wyboru jednego podziału może być osiągnięte przez skalaryzację problemu wielokryterialnego z wektorową funkcją (17) oraz ograniczeniami (4)-(6).

Jedną z koncepcji jest skalaryzacja przez ustalenie ścisłej hierarchii graczy (wygranych graczy) przez zastosowanie wartości Shapleya [1]. Rozwiązanie to jednak silnie promuje gracza, którego cechuje najwyższa wartość Shapleya. Również nie jest jednoznaczne rozwiązanie w sytuacji, gdy wartość Shapleya jest taka sama dla kilku graczy.

Konsekwencją nadania znaczenia graczom jest wykorzystanie skalaryzacji przez zastosowanie sumy ważonej. Wagi mogą być ustalone przez wartość Shapleya lub w inny sposób, podobnie jak w koncepcji skalaryzacji przez ustalenie ścisłej hierarchii graczy. W tym celu można wykorzystać np. inne

przedstawione klasyczne wartości określające pojedynczy podział. Jednak postępowanie takie jest zasadne w sytuacji, gdy taki podział jest poza rdzeniem gry, w przeciwnym wypadku powstaje pytanie o celowość rozważania problemu wielokryterialnego, jeśli znany i akceptowany jest *explicite* pojedynczy podział.

Naturalną konsekwencją stosowania przedstawionych powyżej koncepcji wspomagania wyznaczenia podziału spełniającego postulat racjonalności w sytuacji, gdy niedopuszczalny jest podział ustalony z góry, jest wykorzystanie programowania celowego. Koncepcja ta polega, ogólnie rzecz ujmując, na minimalizacji utworzonej funkcji strat, wynikających z niemożliwości osiągnięcia z góry założonego celu. Funkcja strat może mieć następującą postać

$$\max_i |x_i^* - x_i| \rightarrow \min \quad (18)$$

gdzie x_i^* oznacza ustaloną ad hoc wygraną i -tego gracza. Rozwiązanie programu z funkcją celu (18) oraz ograniczeniami dotyczącymi racjonalności (4)-(6), nie prowadzi jednak w ogólnym przypadku do wyznaczenia jednej imputacji, więc bez dodatkowych heurystyk nie można jednoznacznie rozwiązać zagadnienia. Koncepcja ta jednak może mieć charakter użyteczny w procesie wspomagania podejmowania decyzji dotyczącej podziału wygranej. Zostanie to pokazane w dalszej części pracy przy analizie przykładu.

Koncepcję wykorzystania metody punktu referencyjnego do wspomagania wyboru imputacji zaproponował Wierzbicki w pracy [19]. Metoda punktu referencyjnego pierwotnie została zaproponowana jako metoda wspomagania decyzji wielokryterialnych, w której maksymalizuje się odpowiednio zdefiniowaną funkcję osiągnięcia. Wierzbicki proponuje następującą postać funkcji osiągnięcia

$$\min_{i \in N} \sigma_i(x_i, x_i^r) + \varepsilon \sum_{i=1}^n \sigma_i(x_i, x_i^r) \rightarrow \max \quad (19)$$

gdzie ε jest dodatnim parametrem umożliwiającym jednoznaczne wskazanie imputacji, $\sigma_i(x_i, x_i^r)$ jest częściową funkcją osiągnięcia dla i -tego gracza, kreśloną w następujący sposób

$$\sigma_i(x_i, x_i^r) = \begin{cases} \frac{(x_i - x_i^l)}{(x_i^r - x_i^l)} & \text{dla } x_i \in [x_i^l, x_i^r] \\ 1 + \alpha \frac{(x_i - x_i^r)}{(x_i^u - x_i^r)} & \text{dla } x_i \in [x_i^r, x_i^u] \end{cases} \quad (20)$$

gdzie x_i^r oznacza wartość postulowaną (referencyjną) wypłaty i -tego gracza, α jest dodatnim parametrem optymalizacyjnym, zapewniającym wklęsłość funkcji osiągnięcia.

Opisane wyżej metody zostaną zaprezentowane na przykładzie gry trzech przedsiębiorstw o rynek.

3. Analiza przykładu obliczeniowego – gra przedsiębiorstw o rynek

Rozpatrywana jest hipotetyczna, aczkolwiek możliwa, sytuacja, którą można przedstawić w postaci gry kooperacyjnej z koalicjami.

Na lokalnym rynku działają trzy przedsiębiorstwa: A, B i C, dzieląc między sobą rynek. Dotychczasowa działalność przedsiębiorstw jest jednak zagrożona przez silnego konkurenta spoza rozpatrywanego rynku, który podjął intensywne działania marketingowe w celu wejścia na ten rynek.

Każde z rozpatrywanych przedsiębiorstw dysponuje innymi zasobami. Przedsiębiorstwa A i B są nie są w stanie funkcjonować samodzielnie po wejściu konkurenta na rynek. Przedsiębiorstwo C samodzielnie utrzyma się na rynku (szacuje się, że zachowa 31%), jednak przy wejściu w koalicję z dowolnym z pozostałych przedsiębiorstw razem obronią większą część rynku (z przedsiębiorstwem A około 50% rynku, a z przedsiębiorstwem B 53%). Podobnie współpraca przedsiębiorstwa A i B pozwoli zachować im większą część rynku (około 60%). Jednak w sytuacji, gdy wszystkie trzy będą kooperować, to obronią rynek przed wejściem na niego zagrażającego im konkurenta.

Niech przedsiębiorstwa A, B i C mają odpowiednio numery 1, 2 i 3. Zbiór numerów graczy $N = \{1, 2, 3\}$, a wartości funkcji charakterystycznej przedstawiają się następująco

$$\begin{aligned}
 v(\emptyset) &= 0 \\
 v(\{1\}) &= 0 \\
 v(\{2\}) &= 0 \\
 v(\{3\}) &= 31 \\
 v(\{1,2\}) &= 60 \\
 v(\{1,3\}) &= 50 \\
 v(\{2,3\}) &= 53 \\
 v(\{1,2,3\}) &= 100
 \end{aligned}
 \tag{21}$$

Władze przedsiębiorstw uzgodniły, że zawiążą jedną wielką koalicję, a podział rynku nastąpi według reguły średniego wkładu każdego z przedsiębiorstw w koalicję. Ponadto ustalono, że kolejność tworzenia koalicji nie jest

istotna i mogła przebiegać w dowolny sposób, w związku z tym ustalono równoprawdopodobne scenariusze kolejności wchodzenia do koalicji. Ustalono więc podział według wartości Shapleya (9), która dla poszczególnych przedsiębiorstw wynosi

$$\begin{aligned}\varphi_1 &= 28,83 \\ \varphi_2 &= 30,33 \\ \varphi_3 &= 40,83\end{aligned}\tag{22}$$

Można zauważyć, że ustalony podział nie jest racjonalny z punktu widzenia racjonalności koalicyjnej, ponieważ dla przedsiębiorstw o nr 1 i 2, bardziej opłacalne jest zawiązanie wspólnej koalicji dwuosobowej – rozwiązanie jest poza rdzeniem

$$v(\{1,2\}) > \varphi_1 + \varphi_2\tag{23}$$

$$60 > 28,83 + 30,33 = 59,16\tag{24}$$

Wobec stwierdzenia powyższego faktu przez przedsiębiorstwa, zdecydowano się na przygotowanie scenariuszy podziału rynku według znanych koncepcji.

Przygotowano zestawienie zawierające zaprezentowane koncepcje klasyczne i przedstawiono w tabeli 1.

Tabela 1

Zestawienie scenariuszy podziału wygranej według rozważanych klasycznych metod

Numer gracza	1	2	3
Wartość Shapleya	28,83	30,33	40,83
Wartość Banzhafa	31,50	33,00	43,50
Normalizowana wartość Banzhafa	29,17	30,56	40,28
Wartość Raiffa	30,59	32,55	36,86
Nucleolus	30,75	33,75	35,50
Punkt Gatelý'ego	30,59	32,55	36,86

Można zauważyć, że podział wyznaczony przez unormowaną wartość Banzhafa (11) również nie jest racjonalny (koalicyjnie) dla koalicji przedsiębiorstw o numerach 1 oraz 2

$$v(\{1,2\}) > \beta_1^* + \beta_2^*\tag{25}$$

$$60 > 29,17 + 30,56 = 59,73\tag{26}$$

Przygotowano również zestawienie prezentujące scenariusze podziału przy wykorzystaniu metod quasi-klasycznych i przedstawiono w tabeli 2.

Tabela 2

Zestawienie scenariuszy podziału wygranej według rozważanych quasi-klasycznych metod

Numer gracza	1	2	3
Programowanie leksykograficzne:			
hierarchia ustalona według dowolnej z przedstawionych klasycznych koncepcji	10,00	50,00	40,00
hierarchia ustalona według numerów graczy:			
1, 2, 3	47,00	22,00	31,00
1, 3, 2	47,00	13,00	40,00
2, 3, 1	10,00	50,00	40,00
2, 1, 3	19,00	50,00	31,00
3, 1, 2	47,00	13,00	40,00
Skalaryzacja przez sumę ważoną przez:			
dowolną z wartości klasycznych	10,00	50,00	40,00
Metoda punktu referencyjnego:			
z wartością Shapleya	29,67	30,33	40,00
z wartością Banzhafa	31,50	28,50	40,00
z nukleolusem	30,75	33,75	35,50
z punktem Gately'ego (wartością Raiffa)	30,59	32,55	36,86

Scenariusze podziału wygranej przy zastosowaniu programowania leksykograficznego uwypuklają znaczenie gracza, który w hierarchii jest najważniejszy, dlatego rozwiązania dla graczy najważniejszych to górne ograniczenia rdzenia – maksymalne możliwe racjonalne wygrane

$$\begin{aligned}
 x_1^u &= 47 \\
 x_2^u &= 50 \\
 x_3^u &= 40
 \end{aligned}
 \tag{27}$$

W koncepcji wykorzystania programowania leksykograficznego sposób podziału wygranej można scharakteryzować w ten sposób, że gracz najważniejszy „zabiera” z wygranej całej koalicji ile może najwięcej i resztę zostawia pozostałym. Dalej kolejny w hierarchii „zabiera” tyle, ile racjonalny podział umożliwia, i tak dalej. W końcu najmniej ważny gracz otrzymuje tyle, ile mu pozostali zostawia.

Koncepcja wykorzystania skalaryzacji przez sumę ważoną implikuje podobne rozwiązania jak w przypadku wcześniejszej koncepcji, z tą różnicą, że mogą wystąpić równe wagi graczy (wypłat graczy), wtedy otrzymane rozwiązanie nie jest zbiorem jednoelementowym. Poza tym metoda ta również uwypukla znaczenie gracza, któremu nadana została największa waga. W istocie swej koncepcja ta realizuje logikę leksykograficzną, gdzie hierarchia graczy jest określona przez wagi.

Metoda punktu referencyjnego jest atrakcyjna, ponieważ zawsze implikuje jedną imputację, która należy do rdzenia gry. W sytuacji, gdy punkt referencyjny nie jest racjonalny, czyli nie należy do rdzenia gry (w analizowanym przykładzie jest to wartość Shapleya i wartość Banzhafa), wtedy otrzymane rozwiązanie można interpretować jako taki podział, który należy do rdzenia i jest najbliższy punktowi referencyjnemu (założonemu rozwiązaniu ad hoc).

Metoda punktu referencyjnego odnosi się do celu, z tą różnicą w odniesieniu do logiki programowania celowego, że maksymalizuje się funkcję osiągnięcia celu (punktu referencyjnego), a w koncepcji programowania celowego minimalizuje się funkcję strat, wynikających z nieosiągnięcia celu.

W związku z powyższym, podobnie można interpretować wynik otrzymany w przypadku wykorzystania logiki programowania celowego, z funkcją start zdefiniowanej przez wyrażenie (18). Z tą jednak różnicą, że w ogólnym przypadku rozwiązaniem nie jest jedna imputacja, lecz zbiór wszystkich imputacji, dla których funkcja strat osiąga minimum. Określenie jednego podziału wymaga zastosowania innych heurystyk – na przykład związanych z hierarchią graczy ustaloną przez wartości będące celem lub w inny sposób, czyli można wykorzystać koncepcję programowania leksykograficznego lub skalaryzacji przez sumę ważoną.

Minimalne z maksymalnych odchyłeń wynosi 0,83, przy przyjęciu wartości Shapleya, jako celu w analizowanym przypadku. Natomiast przy przyjęciu jako celu wartości Banzhafa, minimalne odchylenie wynosi 3,5. Możliwe podziały ze względu na znaczenie graczy, przy celu będącym wartością Shapleya lub wartością Banzhafa, przedstawia tabela 3.

Tabela 3

Zestawienie scenariuszy podziału wygranej przy wykorzystaniu logiki programowania celowego

Numer gracza	1	2	3
Cel: wartość Shapleya, hierarchia: 321 lub 231 lub 213	28,83	31,17	40,00
Cel: wartość Shapleya, hierarchia: 312 lub 132 lub 123	29,67	30,33	40,00
Cel: wartość Banzhafa, hierarchia: 321 lub 231 lub 213	28,00	32,00	40,00
Cel: wartość Banzhafa, hierarchia: 312 lub 132 lub 123	30,50	29,50	40,00

Z powyższego zestawienia na rekomendację, w rozpatrywanym przykładzie, zasługuje podział wynikający z hierarchii ustalonej przez wartość celu – w przypadku obu celów, będzie to hierarchia 321, czyli najważniejszy jest gracz trzeci, później drugi. Pierwszy gracz otrzymuje to, co pozostanie po przydzieleniu wygranej pozostałym koalicjantom.

W przedstawionym przykładzie przedsiębiorstwa wybierając ad hoc sposób implikujący nieracjonalny podział, mogą rozważyć zmianę wcześniejszych ustaleń i wybrać podział wynikający z nukleolusa lub punkt Gately'ego

lub wykorzystać rozwiązania będące najbliższym postulowanemu podziału, będące jednocześnie racjonalne, przez wykorzystanie metody Wierzbickiego albo programowania celowego. Rozwiązania uzyskane przez wykorzystanie programowania leksykograficznego lub skalaryzacji przez sumę ważoną różnicują znaczenie przedsiębiorstw, dlatego w rozważanym przykładzie są mało atrakcyjne.

Podsumowanie

W pracy przedstawiono klasyczne (normatywne) metody podziału wygranej między graczy w grach z koalicjami, których przyjęcie (w szczególności wartości Shapleya lub wartości Banzhafa) może być nieracjonalne. W takich sytuacjach można wspomagać podział wygranej wykorzystując metody quasi-klasyczne, oparte na logice programowania wielokryterialnego, które umożliwiają wskazanie podziału racjonalnego, który będzie możliwie blisko postulowanego (ale nieracjonalnego) podziału. Znaczenie słowa „blisko” jest kluczowe, ponieważ w świetle przedstawionych koncepcji może być rozumiane w różnoraki sposób. Sposób ten zależy od przyjętej metody. Na uwagę zasługuje metoda punktu referencyjnego, która implikuje jednoznaczne rozwiązanie oraz metoda programowania celowego zawężająca rdzeń gry do imputacji spełniające minimum funkcji strat (wynikających z nieosiągnięcia celu), jednak wsparta koncepcją programowania leksykograficznego lub inną racjonalnie uzasadnioną heurystyką również wskazuje pojedynczą, racjonalną imputację.

Koncepcja wykorzystania programowania leksykograficznego oraz skalaryzacji zagadnienia przez zastosowanie sumy ważonej uwypuklają znaczenie gracza najważniejszego. Ponadto, przewiduje się hierarchię graczy, co jest ograniczeniem, ponieważ gracze muszą zaakceptować ten fakt lub sytuacja decyzyjna wskazuje, że taka struktura jest uzasadniona.

Literatura

1. Ameljańczyk A. (1984). Optymalizacja wielokryterialna w problemach sterowania i zarządzania. Ossolineum, Wrocław.
2. Banzhaf J. (1965). Weighted Voting Doesn't Work: A Mathematical Analysis. Rutgers Law Review, 19, 317-343.
3. Bondareva O.N. (1963). Some Applications of Linear Programming Methods to the Theory of Cooperative Games. Problemy Kibernetiki, 10, 119-139.

4. Bronisz P., Kruś L. (2000). Cooperative Game Solution Concepts to Cost Allocation Problem. *European Journal of Operational Research*, 122, 258-271.
5. Gately D. (1974). Sharing the Gains from Customs Unions Among Less Developed Countries: A Game Theoretic Approach. *Journal of Development Economics*, 1, 213-233.
6. Gately D. (1974). Sharing the Gains from Regional Cooperation: A Game Theoretic Application to Planning Investment in Electric Power. *International Economic Review*, 15, 195-208.
7. Grotte J.H. (1971/72). Observations on the Nucleolus and the Central Game. *International Journal of Game Theory*, 1, 173-177.
8. Jakowska-Suwalska K., Wolny M. (2008). Programowanie leksykograficzne w grach kooperacyjnych. W: *Badania operacyjne i systemowe: decyzje, gospodarka, kapitał ludzki i jakość*. Red. J.W. Owsiniński, Z. Nahorski, T. Szapiro. PAN IBS, 64, 65-72.
9. Kalai E., Smorodinsky M. (1975). Other Solutions to Nash's Bargaining Problem. *Econometrica* 43, 513-518.
10. Konarzewska-Gubała E. (1980) Programowanie przy wielorakości celów. PWN, Warszawa.
11. Littlechild S.C., Vaidya K.G. (1976). The Propensity to Disrupt and the Disruption Nucleolus of a Characteristic Function Game. *International Journal of Game Theory*, 5, 151-161.
12. Raiffa H. (1980). *The Art and Science of Negotiations*. Harvard University Press, Cambridge.
13. Schmeidler D. (1969). The Nucleolus of a Characteristic Function Game. *SIAM Journal on Applied Mathematics* 17, 1163-1170.
14. Seo F., Sakawa M. (1987). *Multiple Criteria Decision Analysis in Regional Planning*. Reidel, Dordrecht.
15. Shapley L.S. (1953). A Value for N-person Games. *Annales of Mathematical Studies*, 28, 307-318.
16. Shapley L.S. (1967). On balanced Sets and Cores, *Naval Research Logistics Quarterly*, 14, 453-460.
17. Straffin P. (2001). *Teoria gier*. Wydawnictwo Naukowe Scholar, Warszawa.
18. Wierzbicki A.P. (2005). A Reference Point Approach to Coalition Games. *Journal of Multi-Criteria Decision Analysis*, 13, 81-89.
19. Young H.P., Okada N., Hashimoto T. (1980). *Cost Allocation in Water Resources Development. A Case Study of Sweden*. IIASA Research Report, Laxenburg.

Irena Woroniecka-Leciejewicz

Instytut Badań Systemowych PAN w Warszawie

Wyższa Szkoła Informatyki Stosowanej i Zarządzania w Warszawie

RÓWNOWAGA W GRZE FISKALNO-MONETARNEJ A PRIORYTETY BANKU CENTRALNEGO I RZĄDU

Wstęp

Praca dotyczy zastosowania teorii gier w analizie polityki makroekonomicznej i wyborze policy mix, będącej kombinacją polityki monetarnej i fiskalnej. Celem pracy jest analiza interakcji decyzyjnych i wzajemnych uwarunkowań między władzami monetarnymi i fiskalnymi z wykorzystaniem dwuosobowej gry między bankiem centralnym a rządem, zwanej grą fiskalno-monetarną, ze skończoną liczbą strategii w zakresie polityki pieniężnej i budżetowej, różniących się stopniem ich restrykcyjności/ekspansywności. Przeprowadzono analizę stanów równowagi i Pareto-optymalności rozwiązań w przypadku alternatywnych założeń dotyczących uwarunkowań koniunktury gospodarczej i skuteczności polityki monetarnej i fiskalnej, a także z uwzględnieniem różnych priorytetów banku centralnego i rządu w kształtowaniu polityki makroekonomicznej. Odmienne założenia znajdują odzwierciedlenie w inaczej zdefiniowanych wypłatach w grze, a w konsekwencji w stanach równowagi gry i ich Pareto-optymalności. Początkowe założenie, że bank centralny dąży do minimalizacji inflacji, a rząd do maksymalizacji realnego wzrostu gospodarczego w dalszej części skorygowano przyjmując, że celem władz monetarnych jest osiągnięcie pożądanego poziomu inflacji, tzw. celu inflacyjnego, podczas gdy władze fiskalne dążą do osiągnięcia pożądanego (zaplanowanego) wzrostu gospodarczego.

Wśród ważnych prac z tej dziedziny należy wskazać na prace takich autorów jak: Binder [3], Bennett i Loayza [1], Wojtyna [17; 18], Marszałek [9], Sargent i Wallace [13], Eijffinger, DeHaan [5], Blackburn, Christensen [2], Walsh [16], Gjedrem [6]. Mają oni dorobek bądź w dziedzinach związanych ściśle z analizą współzależności między polityką pieniężną i fiskalną, między

innymi w takich zagadnieniach, jak problem niezależności banku centralnego, znaczenie międzyokresowego ograniczenia budżetowego i oczekiwań inflacyjnych, tzw. nieprzyjemna arytmetyka monetarystyczna (unpleasant monetarist arithmetic), znaczenie aspektów jakościowych, przede wszystkim wiarygodności i przejrzystości działań banku centralnego i rządu, czy tzw. problem jednorękiego decydenta (one-armed policymaker), związany z niewystarczającą skutecznością instrumentów polityki monetarnej w niesprzyjających warunkach ekonomicznych (wysoki dług publiczny) bądź w zastosowaniu teorii gier w badaniu tych interakcji. Praca stanowi kontynuację prowadzonych badań [20; 21; 22; 23].

1. Koncepcja gry monetarno-fiskalnej

Analizowana w niniejszej pracy gra fiskalno-monetarna jest dwuosobową grą między bankiem centralnym i rządem. Jest to jednoetapowa gra o sumie niezerowej z pełną informacją. Każdy z graczy podejmuje decyzje samodzielnie, biorąc pod uwagę prawdopodobną reakcję drugiego gracza. Strategie rządowe oznaczają strategie polityki fiskalnej, różniące się stopniem restrykcyjności polityki. Jako miernik stopnia restrykcyjności polityki fiskalnej przyjęto poziom deficytu budżetowego w relacji do PKB. Strategie banku centralnego oznaczają różniące się stopniem restrykcyjności strategie polityki monetarnej, przy czym jako wyznacznik restrykcyjności polityki pieniężnej przyjęto wysokość realnej stopy procentowej. Wypłatą banku centralnego jest poziom inflacji, zakłada się, że władze monetarne dążą do uzyskania jak najniższego jej poziomu. Wypłatą, która jest maksymalizowana przez władze fiskalne (rząd) jest tempo wzrostu PKB.

Tabela 1

Gra ze skończoną liczbą strategii fiskalnych i monetarnych

Tablica wypłat		Bank centralny			
		Strategia monetarna $M1$ (realna stopa procentowa r_1)	Strategia monetarna $M2$ (realna stopa procentowa r_2)	...	Strategia monetarna Mn (realna stopa procentowa r_n)
Rząd	Strategia fiskalna $F1$ (deficyt budżetowy b_1)	p_{11} y_{11}	p_{12} y_{12}	...	p_{1n} y_{1n}
	Strategia fiskalna $F2$ (deficyt budżetowy b_2)	p_{21} y_{21}	p_{22} y_{22}	...	p_{2n} y_{2n}
	...			...	
	Strategia fiskalna Fm (deficyt budżetowy b_m)	p_{m1} y_{m1}	p_{m2} y_{m2}	...	p_{mn} y_{mn}

Tabela 1 przedstawia tablicę wypłat dla tak zdefiniowanej gry. Wypłaty zostały oznaczone w następujący sposób: y_{ij} – wypłata rządu (tempo wzrostu PKB) w przypadku, gdy rząd stosuje strategię fiskalną F_i , a bank centralny strategię monetarną M_j , p_{ij} – wypłata banku centralnego (inflacja) w tej samej sytuacji strategicznej. Symbolem r_j oznaczono realną stopą procentową przypisaną j -tej strategii pieniężnej, natomiast symbolem b_i – deficyt budżetowy w relacji do PKB, charakteryzujący i -tą strategię fiskalną. Koncepcja powyższej gry dla dwóch jakościowo różnych strategii fiskalnych i monetarnych: polityki restrykcyjnej i ekspansywnej, przy założeniu liniowych zależności między stanem gospodarki (charakteryzowanym przez wzrost gospodarczy i inflację) a instrumentami polityki makroekonomicznej (deficytem budżetowym i stopą procentową) została przedstawiona w pracach [20; 21; 22; 23].

2. Równowaga w grze monetarno-fiskalnej z dwoma strategiami dla nieliniowych zależności

W niniejszym rozdziale przedstawiony zostanie przykład gry uwzględniającej dwie jakościowo różne strategie: restrykcyjną i ekspansywną zarówno po stronie polityki fiskalnej jak i monetarnej (tabela 2), przy odejściu od założeń o liniowym charakterze zależności między stanem gospodarki a stosowanymi instrumentami polityki makroekonomicznej. Przyjmuje się, że bank centralny, dążąc do obniżania inflacji (p), wybiera między polityką bardziej restrykcyjną, charakteryzującą się wyższą realną stopą procentową (r) a polityką mniej restrykcyjną. Władze fiskalne skłaniają się bądź do wyboru polityki bardziej restrykcyjnej, której towarzyszy niższy poziom deficytu budżetu państwa, bądź bardziej ekspansywnej (wyższy deficyt), dążąc do osiągnięcia jak najwyższego wzrostu PKB. Lewa kolumna odzwierciedla restrykcyjną politykę pieniężną, prawa zaś ekspansywną, analogicznie górny wiersz oznacza restrykcyjną, a dolny ekspansywną politykę fiskalną.

Gra analizowana jest przy przyjęciu alternatywnych założeń dotyczących wpływu instrumentów polityki fiskalnej i monetarnej na stan gospodarki, odzwierciedlany przez wzrost PKB i inflację. Rozpatrywane są dwa warianty. W obu przyjmuje się założenie, że wzrost realnej stopy procentowej, *ceteris paribus*, wywołuje spadek tempa wzrostu PKB oraz ograniczenie inflacji, a wzrost deficytu budżetowego, *ceteris paribus*, przyczynia się do wzrostu inflacji. Różnica dotyczy wpływu deficytu budżetowego na realny wzrost produkcji w gospodarce. W wariantcie A zakłada się, że wzrost deficytu budżetu państwa, *ceteris paribus*, powoduje zwiększenie tempa wzrostu PKB, podczas, gdy w wariantcie B – ogranicza tempo wzrostu produkcji.

Tabela 2

Gra monetarno-fiskalna z dwoma strategiami

Tablica wypłat			Bank centralny	
			Strategia M1	Strategia M2
			r_1	r_2
Rząd	Strategia F1	b_1	p_{11} y_{11}	p_{12} y_{12}
	Strategia F2	b_2	p_{21} y_{21}	p_{22} y_{22}

Tabela 3 przedstawia tablicę wypłat w grze dla wariantu A. Formuły określające zależność tempa wzrostu PKB oraz inflacji od zmian deficytu budżetowego i stopy procentowej wyprowadzono wykorzystując rozwinięcie w szereg Taylora z dokładnością do drugiej pochodnej

$$f(x + \Delta x, z + \Delta z) = f(x, z) + \frac{\partial f}{\partial x} \Delta x + \frac{\partial f}{\partial z} \Delta z + \frac{1}{2} \left(\frac{\partial^2 f}{\partial x^2} \Delta x^2 + 2 \frac{\partial^2 f}{\partial x \partial z} \Delta x \Delta z + \frac{\partial^2 f}{\partial z^2} \Delta z^2 \right) \quad (1)$$

Najniższa inflacja i jednocześnie najniższy wzrost gospodarczy występuje w przypadku wyboru kombinacji restrykcyjnych polityk: monetarnej i fiskalnej (lewy górny róg tablicy wypłat). Wraz z obniżaniem stopy procentowej (przejście do prawej kolumny) zwiększa się inflacja i rośnie tempo wzrostu PKB

$$p + \frac{\partial p}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial r^2} \Delta r^2, \quad y + \frac{\partial y}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial r^2} \Delta r^2 \quad (2)$$

przy czym: $\Delta b = b_2 - b_1 > 0$, $\Delta r = r_2 - r_1 < 0$

Również na skutek rosnącego deficytu budżetowego następuje wzrost inflacji i produkcji (przejście do dolnego wiersza)

$$p + \frac{\partial p}{\partial b} \Delta b + \frac{1}{2} \frac{\partial^2 p}{\partial b^2} \Delta b^2, \quad y + \frac{\partial y}{\partial b} \Delta b + \frac{1}{2} \frac{\partial^2 y}{\partial b^2} \Delta b^2 \quad (3)$$

Tabela 3

Równowaga w grze monetarno-fiskalnej z dwoma strategiami dla nieliniowych zależności. Wariant A

Wypląty		Bank centralny	
		wyższa stopa procentowa r_1	niższa stopa procentowa r_2
Rząd	niższy deficyt b_1	p y	$p + \frac{\partial p}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial r^2} \Delta r^2$ $y + \frac{\partial y}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial r^2} \Delta r^2$
	wyższy deficyt b_2	$p + \frac{\partial p}{\partial b} \Delta b + \frac{1}{2} \frac{\partial^2 p}{\partial b^2} \Delta b^2$ $y + \frac{\partial y}{\partial b} \Delta b + \frac{1}{2} \frac{\partial^2 y}{\partial b^2} \Delta b^2$	$p + \frac{\partial p}{\partial b} \Delta b + \frac{\partial p}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial b^2} \Delta b^2 + \frac{\partial^2 p}{\partial b \partial r} \Delta b \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial r^2} \Delta r^2$ $y + \frac{\partial y}{\partial b} \Delta b + \frac{\partial y}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial b^2} \Delta b^2 + \frac{\partial^2 y}{\partial b \partial r} \Delta b \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial r^2} \Delta r^2$

Najwyższą inflacją, ale i najszybszym wzrostem produkcji charakteryzuje się gospodarka, gdy zarówno polityka pieniężna, jak i fiskalna mają charakter ekspansywny (prawy dolny róg tabeli)

$$p + \frac{\partial p}{\partial b} \Delta b + \frac{\partial p}{\partial r} \Delta r + \frac{1}{2} \left(\frac{\partial^2 p}{\partial b^2} \Delta b^2 + 2 \frac{\partial^2 p}{\partial b \partial r} \Delta b \Delta r + \frac{\partial^2 p}{\partial r^2} \Delta r^2 \right) \quad (4)$$

Optymalną strategią banku centralnego w przypadku zastosowania przez rząd restrykcyjnej polityki fiskalnej (niższy deficyt b_1) jest wybór restrykcyjnej polityki monetarnej (wyższa stopa procentowa r_1), ponieważ wybierana jest niższa inflacja w pierwszym wierszu ($p_{11} < p_{12}$). W przypadku wyboru przez władze fiskalne ekspansywnej polityki budżetowej, optymalną strategią władz monetarnych pozostaje polityka restrykcyjna, bowiem w drugim wierszu także wybierany jest niższy poziom inflacji ($p_{21} < p_{22}$). Stąd wniosek, że restrykcyjna polityka pieniężna stanowi dla banku centralnego strategię dominującą, która jest optymalna niezależnie od tego, jaką strategię polityki fiskalnej wybierze rząd – restrykcyjną czy ekspansywną.

Przeprowadzając analogiczne rozumowanie dla alternatywnych decyzji władz fiskalnych, które dążą do zwiększania realnego wzrostu PKB, dochodzimy do konkluzji, że optymalną strategią budżetową w odpowiedzi na restrykcyjną politykę pieniężną jest polityka ekspansywna, ponieważ wybierane jest wyższe tempo wzrostu PKB w pierwszej kolumnie ($y_{21} > y_{11}$), ale również w przeciwnym wypadku, gdy bank centralny realizuje restrykcyjną politykę monetarną, optymalną strategią fiskalną także jest polityka o charakterze ekspansywnym – odpowiada temu wybór wyższego wzrostu PKB w drugiej kolumnie ($y_{22} > y_{12}$). Wynika z tego, że rząd, podobnie jak bank centralny, ma strategię dominującą, którą jest ekspansywna polityka fiskalna, jest ona bowiem z punktu widzenia rządu strategią optymalną niezależnie od tego, jakie decyzje o wysokości stóp procentowych podejmie bank centralny.

Stan równowagi znajduje się w lewym dolnym rogu tablicy wypłat. Jest to nie tylko równowaga Nasha, jest to równowaga determinowana przez strategię dominującą. Prowadzi do wyboru przez podmioty odpowiedzialne za politykę makroekonomiczną restrykcyjnej polityki monetarnej z jednej strony i ekspansywnej polityki budżetowej, z drugiej.

3. Problem Pareto-optymalności rozwiązań w grze fiskalno-monetarnej z dwoma startegiami i jego interpretacja

Interesujące jest pytanie, czy równowaga, wyznaczana przez strategię dominującą władz fiskalnych i monetarnych, jest Pareto-optymalna. W tym sensie szczególnie ważne jest porównanie dwóch wariantów rozwiązań strategicznych: stanu równowagi, odzwierciedlającego połączenie ekspansywnej polityki fiskalnej i restrykcyjnej monetarnej (lewy dolny róg tablicy wypłat) oraz alternatywnego rozwiązania – kombinacji restrykcyjnej polityki budżetowej i ekspansywnej pieniężnej (prawy górny róg tablicy wypłat).

Ponieważ, jak się wydaje, nie ma odpowiedzi uniwersalnej, rozważone zostaną cztery możliwe przypadki

Przypadek A1

$$\begin{aligned} \frac{\partial y}{\partial b} \Delta b + \frac{\partial^2 y}{\partial b^2} \Delta b^2 &> \frac{\partial y}{\partial r} \Delta r + \frac{\partial^2 y}{\partial r^2} \Delta r^2 \wedge \\ \frac{\partial p}{\partial b} \Delta b + \frac{\partial^2 p}{\partial b^2} \Delta b^2 &< \frac{\partial p}{\partial r} \Delta r + \frac{\partial^2 p}{\partial r^2} \Delta r^2 \end{aligned} \quad (5)$$

Warunek ten oznacza, że po pierwsze, zwiększenie tempa wzrostu PKB wywołane wzrostem deficytu budżetowego, ceteris paribus, stanowi efekt silniejszy niż zwiększenie wzrostu PKB wywołane obniżeniem stopy procentowej,

ceteris paribus. Po drugie, efekt inflacyjny rosnącego deficytu finansów publicznych jest silniejszy w porównaniu ze skutkiem inflacyjnym obniżki stopy procentowej. W pozostałych przypadkach interpretacja założeń jest analogiczna i zostanie pominięta.

Przypadek A2

$$\begin{aligned} \frac{\partial y}{\partial b} \Delta b + \frac{\partial^2 y}{\partial b^2} \Delta b^2 &> \frac{\partial y}{\partial r} \Delta r + \frac{\partial^2 y}{\partial r^2} \Delta r^2 \wedge \\ \frac{\partial p}{\partial b} \Delta b + \frac{\partial^2 p}{\partial b^2} \Delta b^2 &> \frac{\partial p}{\partial r} \Delta r + \frac{\partial^2 p}{\partial r^2} \Delta r^2 \end{aligned} \quad (6)$$

Przypadek A3

$$\begin{aligned} \frac{\partial y}{\partial b} \Delta b + \frac{\partial^2 y}{\partial b^2} \Delta b^2 &< \frac{\partial y}{\partial r} \Delta r + \frac{\partial^2 y}{\partial r^2} \Delta r^2 \wedge \\ \frac{\partial p}{\partial b} \Delta b + \frac{\partial^2 p}{\partial b^2} \Delta b^2 &< \frac{\partial p}{\partial r} \Delta r + \frac{\partial^2 p}{\partial r^2} \Delta r^2 \end{aligned} \quad (7)$$

Przypadek A4

$$\begin{aligned} \frac{\partial y}{\partial b} \Delta b + \frac{\partial^2 y}{\partial b^2} \Delta b^2 &< \frac{\partial y}{\partial r} \Delta r + \frac{\partial^2 y}{\partial r^2} \Delta r^2 \wedge \\ \frac{\partial p}{\partial b} \Delta b + \frac{\partial^2 p}{\partial b^2} \Delta b^2 &> \frac{\partial p}{\partial r} \Delta r + \frac{\partial^2 p}{\partial r^2} \Delta r^2 \end{aligned} \quad (8)$$

Tabela 4 przedstawia preferencje w analizowanej dla powyższych czterech przypadków założeń. W przypadku A1 stan równowagi reprezentujący kombinację ekspansywnej polityki fiskalnej i restrykcyjnej monetarnej stanowi rozwiązanie Pareto-optymalne – tempo wzrostu PKB jest wyższe, a inflacja niższa w porównaniu z rozwiązaniem alternatywnym z prawego górnego rogu oznaczającym połączenie restrykcyjnej polityki budżetowej i ekspansywnej pieniężnej. Przyjęcie strategii dominujących jest korzystniejszym rozwiązaniem dla obu graczy: rządu jak i banku centralnego. Również w przypadkach A2 i A3 równowaga jest Pareto-optymalna, przy czym każdy z tych przypadków stanowi zwierciadlane odbicie drugiego. Przy założeniach A2 punkt równowagi jest korzystniejszym rozwiązaniem z punktu widzenia władz fiskalnych, pozwala osiągnąć wyższy wzrost gospodarczy, jednak gorszym w ocenie banku centralnego, ponieważ ten pozytywny efekt przyśpieszenia wzrostu obciążony jest kosztem inflacyjnym. W przypadku A3 stan równowagi jest korzystniejszy z punktu widzenia władz monetarnych (niższa inflacja), jednak mniej korzystny dla władz fiskalnych (niższy wzrost PKB) w porównaniu z kombinacją restrykcyjnej polityki budżetowej i ekspansywnej pieniężnej.

Tabela 4

Preferencje w grze monetarno-fiskalnej. Przypadki A1-A4

		Bank centralny	
		r_1	r_2
Rząd	B_1	1 / 4	3 / 3
	B_2	2 / 2	4 / 1

		Bank centralny	
		r_1	r_2
Rząd	b_1	1 / 4	2 / 3
	b_2	3 / 2	4 / 1

		Bank centralny	
		r_1	r_2
Rząd	b_1	1 / 4	3 / 2
	b_2	2 / 3	4 / 1

Przypadek A4 – dylemat więźnia

		Bank centralny	
		r_1	r_2
Rząd	b_1	1 / 4	2 / 2
	b_2	3 / 3	4 / 1

W przypadku A4 stan równowagi charakteryzuje się gorszymi, w porównaniu z alternatywną parą strategii, wskaźnikami gospodarczymi (niższy wzrost PKB i wyższa inflacja), co oznacza, że wybór strategii dominujących nie stanowi rozwiązania Pareto-optimalnego. Przypadek ten odzwierciedla sytuację znaną w literaturze jako dylemat więźnia, kiedy występuje konflikt między racjonalnością indywidualną reprezentowaną przez kryterium strategii dominującej a racjonalnością grupową reprezentowaną przez kryterium Pareto-optimalności. Sytuacja ta wskazuje na celowość koordynacji polityki banku centralnego i rządu.

Na podstawie powyższej analizy można sformułować wniosek dotyczący czynników, od których zależy, czy równowaga Nasha (w analizowanym problemie wyznaczana przez strategię dominującą) stanowi jednocześnie rozwiązanie Pareto-optimalne czy nie. W szerszym kontekście pytanie to dotyczy istotnego i często podejmowanego w dyskusjach ekonomicznych problemu – czy niezależne podejmowanie decyzji przez władze fiskalne i monetarne prowadzi do rozwiązań optymalnych w sensie Pareto czy też nieoptymalnych, i w związku z tym wymaga koordynacji działań. Uzyskane wyniki wskazują, że Pareto-optimalność rozwiązań zależy od skuteczności instrumentów polityki monetarnej i fiskalnej w oddziaływaniu na gospodarkę, w szczególności na wzrost gospodarczy i na inflację.

4. Równowaga i Pareto-optymalność rozwiązań gry fiskalno-monetarnej przy alternatywnych założeniach

Dotychczas gra między podmiotami odpowiedzialnymi za politykę monetarną i fiskalną analizowana była przy przyjęciu założeń zgodnie z wariantem A, który w krótkim okresie wydaje się bardziej realistycznie odzwierciedlać badane interakcje między polityką pieniężną a budżetową. Można jednak rozważyć również alternatywne założenia (wariant B), przyjmując, w przeciwieństwie do wariantu A, że wzrost deficytu budżetowego, ceteris paribus, wywołuje ograniczenie tempa wzrostu PKB.

Tabela 5

Równowaga w grze monetarno-fiskalnej z dwoma strategiami dla nieliniowych zależności. Wariant B

Wyплаты		Bank centralny	
		wyższa stopa procentowa r_1	niższa stopa procentowa r_2
Rząd	niższy deficyt b_1	p $y + \frac{\partial y}{\partial b} \Delta b + \frac{1}{2} \frac{\partial^2 y}{\partial b^2} \Delta b^2$	$p + \frac{\partial p}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial r^2} \Delta r^2$ $y + \frac{\partial y}{\partial b} \Delta b + \frac{\partial y}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial b^2} \Delta b^2 + \frac{\partial^2 y}{\partial b \partial r} \Delta b \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial r^2} \Delta r^2$
	wyższy deficyt b_2	$p + \frac{\partial p}{\partial b} \Delta b + \frac{1}{2} \frac{\partial^2 p}{\partial b^2} \Delta b^2$ y	$p + \frac{\partial p}{\partial b} \Delta b + \frac{\partial p}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial b^2} \Delta b^2 + \frac{\partial^2 p}{\partial b \partial r} \Delta b \Delta r + \frac{1}{2} \frac{\partial^2 p}{\partial r^2} \Delta r^2$ $y + \frac{\partial y}{\partial r} \Delta r + \frac{1}{2} \frac{\partial^2 y}{\partial r^2} \Delta r^2$

Tabela 5 przedstawia tablicę wypłat dla założeń zgodnych z wariantem B. Zakłada się przy tym, że: $\Delta b = b_2 - b_1 < 0$, $\Delta r = r_2 - r_1 < 0$. Nastąpiła zmiana dotycząca znaku Δb , ponieważ najniższe tempo wzrostu PKB i najniższa inflacja występuje w tym przypadku (wariant założeń B) w sytuacji połączenia restrykcyjnych rodzajów polityki: zarówno monetarnej, jak i polityki fiskalnej – lewy dolny, a nie jak poprzednio – lewy górny róg tabeli wypłat. Wraz z ograniczaniem *ceteris paribus*, deficytu budżetowego (przejście do górnego wiersza), następuje efekt szybszego wzrostu gospodarczego i jednocześnie hamowania inflacji, natomiast w wyniku obniżania się stopy procentowej (przejście do prawej kolumny) można się spodziewać przyspieszenia wzrostu produkcji, ale i zarazem wyższej inflacji. Najwyższe tempo wzrostu gwarantuje połączenie ekspansywnej polityki pieniężnej i restrykcyjnej fiskalnej (prawy górny róg tabeli), podczas gdy najwyższą inflacją charakteryzuje się gospodarka, gdy obie polityki mają charakter ekspansywny (prawy dolny róg tabeli).

Jak można zauważyć, bank centralny kierując się minimalizacją inflacji będzie, niezależnie od decyzji władz fiskalnych, wybierał restrykcyjną politykę monetarną, która stanowi dla banku centralnego, analogiczne jak w wariantcie A, strategię dominującą. Rząd natomiast, dążąc do maksymalizacji realnego wzrostu PKB, będzie niezależnie od decyzji władz monetarnych, wybierał restrykcyjną politykę fiskalną. W tym przypadku strategią dominującą rządu jest restrykcyjna, nie ekspansywna jak w wariantcie A, polityka fiskalna. Równowaga w powyższej grze (wariant B) jest wyznaczona przez strategię dominującą i prowadzi do wyboru zarówno przez władze monetarne, jak i fiskalne – polityki restrykcyjnej. Stan równowagi usytuowany jest w lewym górnym rogu tablicy.

W wariantcie B równowaga stanowi jednocześnie rozwiązanie Pareto-optymalne, gdyż stan równowagi gwarantuje najniższy poziom inflacji, podczas gdy inne rozwiązania dają wprawdzie możliwość szybszego wzrostu gospodarczego, ale zawsze kosztem wyższej inflacji. Można jednak, analogicznie jak w wariantcie A, prześledzić poszczególne przypadki i na podstawie przyjętych założeń określić dla nich preferencje banku centralnego i rządu. Preferencje te zawarte są w tabeli 6. We wszystkich przypadkach równowaga ma charakter Pareto-optymalny, nie ma bowiem innego rozwiązania wśród trzech pozostałych, które nie pogarszając wyniku w zakresie inflacji poprawiałoby wynik dotyczący wzrostu PKB. Warto zwrócić uwagę, że oprócz stanu równowagi oznaczającego wybór restrykcyjnej polityki tak w sferze budżetowej, jak i pieniężnej, również sięgnięcie po kombinację restrykcyjnej polityki fiskalnej i ekspansywnej monetarnej (prawy górny róg tablicy) jest rozwiązaniem Pareto-optymalnym.

Tabela 6

Preferencje w grze monetarno-fiskalnej. Przypadki B1-B4

		Bank centralny	
		r_1	r_2
Rząd	B_1	1 2	3 1
	B_2	2 4	4 3

		Bank centralny	
		r_1	r_2
Rząd	b_1	1 2	2 1
	b_2	3 4	4 3

		Bank centralny	
		R_1	r_2
Rząd	b_1	1 3	3 1
	b_2	2 4	4 2

		Bank centralny	
		r_1	r_2
Rząd	b_1	1 3	2 1
	b_2	3 4	4 2

5. Równowaga w grze ze skończoną liczbą strategii a priorytety władz monetarnych i fiskalnych

Przy rozważanych dotychczas założeniach, oznaczających między innymi dążenie władz fiskalnych do maksymalizacji wzrostu gospodarczego i minimalizacji inflacji przez bank centralny, równowaga w grze fiskalno-monetarnej ze skończoną liczbą strategii jest determinowana przez strategie dominujące, składające do wyboru bądź restrykcyjnej polityki pieniężnej i ekspansywnej budżetowej (wariant A), bądź obu restrykcyjnych polityk (wariant B).

Wydaje się, że odmienne wyniki można uzyskać odchodząc od założenia, że bank centralny dąży do minimalizacji inflacji na rzecz innego założenia, że celem polityki monetarnej jest osiąganie pożądanego jej poziomu, odzwierciedlającego cel inflacyjny polityki pieniężnej, a także konsekwentnie, zastępując założenie, że władze rządowe dążą do maksymalizacji wzrostu gospodarczego założeniem, że dążą one do osiągania pożądanego, planowanego poziomu tempa wzrostu PKB.

Skoncentrujmy na wstępie uwagę na wariantcie A założeń dotyczących wpływu instrumentów polityki makroekonomicznej na stan gospodarki, który w krótkim okresie wydaje się bardziej realistycznie odzwierciedlać badane interakcje między polityką pieniężną a budżetową.

Tabela 8

[illegible]

Tabela 9

Równowaga w grze fiskalno-monetarnej ze skończoną liczbą strategii. Wariant A

		Polityka monetarna												
		r_1	← restrykcyjna				...	ekspansywna →				r_n		
Polityka fiskalna	b_1					p_e								
	↑ restrykcyjna					p_e								
						p_e								
						p_e								
					p_e									
					p_e									
					p_e									
	...				p_e						y_e	y_e	y_e	y_e
	↓ ekspansywna				p_e			y_e	y_e	y_e	y_e			
					p_e	y_e	y_e	y_e						
				y_e	N									
					p_e									
			y_e	p_e										
				p_e										
b_m	y_e	p_e												

Dla odmiennych założeń dotyczących celów polityki fiskalnej i monetarnej przyjmujących, że władze monetarne i fiskalne dążą do osiągnięcia pożądaných wartości odpowiednio inflacji i tempa wzrostu gospodarczego, stan równowagi nie jest, jak poprzednio, wyznaczany przez strategie dominujące, gdyż takimi nie dysponuje ani rząd, ani bank centralny. Stan równowagi Nasha, oznaczony symbolem N, znajduje się w innym miejscu tabeli wypłat niż stan równowagi wyznaczany przez strategie dominujące przy poprzednich założeniach, oznaczony symbolem D (tabela 9). Oznacza to, że władze dążące do równowagi w grze, zmieniają swoje decyzje w zakresie policy-mix. Następuje przesunięcie polityki fiskalnej ze zdecydowanie ekspansywnej w kierunku polityki o charakterze bardziej neutralnym, aczkolwiek w dalszym ciągu o znacznym stopniu ekspansji. Również w zakresie polityki monetarnej obserwuje się zmianę z polityki ekstremalnie restrykcyjnej w kierunku bardziej umiarkowanej, choć charakteryzującej się dość znacznym stopniem restrykcyjności.

Podsumowanie

W pracy przedstawiono i przeanalizowano grę monetarno-fiskalną jako przykład dwuosobowej gry o sumie niezerowej, w której władze monetarne i fiskalne podejmują decyzje samodzielnie. Gra uwzględnia skończoną liczbę różnych strategii w zakresie polityki fiskalnej i monetarnej: od restrykcyjnej do ekspansywnej. Za miernik stopnia restrykcyjności/ekspansywności polityki fiskalnej przyjęto deficyt budżetowy w relacji do PKB, a polityki pieniężnej – realną stopę procentową. Budując tablicę wypłat gry założono początkowo, że bank centralny dąży do minimalizacji inflacji, a rząd do maksymalizacji realnego wzrostu gospodarczego, a w dalszej części, że obie władze dążą do osiągnięcia pożądanego poziomu inflacji i tempa wzrostu gospodarczego.

Podjęto próbę analizy powyższej gry monetarno-fiskalnej rezygnując z założenia o liniowych zależnościach między miernikami stanu gospodarki (wzrostu gospodarczego i inflacji) a instrumentami polityki mix: deficytem budżetowym w relacji do PKB i stopą procentową. Wykorzystując wzór Taylora, wyprowadzono formuły na wartości występujące w tabeli wypłat w grze dla nieliniowych zależności, a następnie przeprowadzono analizę stanów równowagi gry i ich Pareto-optymalności.

W zależności od przyjętych założeń dotyczących uwarunkowań koniunktury gospodarczej, przede wszystkim wpływu deficytu budżetowego na wzrost PKB, niezależnie działające władze monetarne i fiskalne dążą, zgodnie z równowagą wyznaczaną przez strategię dominującą, do wyboru restrykcyjnej polityki pieniężnej i ekspansywnej budżetowej (wariant A) bądź do wyboru obu restrykcyjnych polityk (wariant B). W wariantcie A rozwiązanie wyznaczone przez stan równowagi nie zawsze jest Pareto-optymalne. W szczególnym przypadku dążenie do strategii dominujących może oznaczać wybór rozwiązania nieoptimalnego w sensie Pareto, analogicznie jak w dylemacie więźnia, wtedy lepszą decyzję jest w stanie zapewnić jedynie koordynacja obu rodzajów polityki. W wariantcie B we wszystkich rozpatrywanych przypadkach równowaga stanowi rozwiązanie Pareto-optymalne.

Analiza gry dla wariantu A przy przyjęciu założenia, że bank centralny dąży do osiągnięcia pożądanego poziomu inflacji, podobnie jak celem władz fiskalnych jest osiągnięcie pożądanego tempa wzrostu gospodarczego wskazuje, że w takim przypadku nie ma strategii dominujących. Inne usytuowanie stanu równowagi Nasha oznacza wybór innej polityki-mix – stosunkowo restrykcyjnej polityki pieniężnej i relatywnie ekspansywnej fiskalnej. W porównaniu z wynikami uzyskanymi dla wcześniejszych założeń nastąpiło przesunięcie z radykalnie restrykcyjnej polityki monetarnej i zdecydowanie ekspansywnej polityki fiskalnej w kierunku polityki o charakterze bardziej neutralnym.

Literatura

1. Bennett, N. Loayza, H. (2001). Policy Biases when the Monetary and Fiscal Authorities have Different Objectives. Central Bank of Chile Working Papers, 66, 299-330.
2. Blackburn K., Christensen M. (1989). Monetary Policy and Policy Credibility: Theories and Evidence. *Journal of Economic Literature*, 27, 1-45.
3. Blinder A. (1983). Issues in the Coordination of Monetary and Fiscal Policy. W: *Monetary Policy in the 1980s*. Federal Reserve Bank of Kansas City, 3-34.
4. Blinder A. (2000). Central Bank Credibility: Why Do We Care? How Do We Build It? *American Economic Review*, 2000 December, 1421-1431.
5. Eijffinger W., DeHaan J. (1996). *The Political Economy of Central Bank Independence*. Princeton University, Princeton.
6. Gjedrem S. (2001). Monetary Policy – the Importance of Credibility and Confidence, *BIS Review*, 7, 1-13.
7. Kokoszcyński R. (2004). *Współczesna polityka pieniężna w Polsce*. PWE, Warszawa.
8. Kot A. (2003). Metody kwantyfikacji restrykcyjności monetarnej, fiskalnej oraz policy mix w krajach akcesyjnych. *Bank i Kredyt*, 6.
9. Marszałek P. (2005). Zastosowanie teorii gier do badania koordynacji polityki pieniężnej i polityki fiskalnej. W: *Studia z bankowości centralnej*. Red. W. Przybylska-Kapuścińska. *Zeszyty Naukowe*. AE, Poznań, 56, 224-247.
10. Nordhaus W.D. (1994). Policy Games: Coordination and Independence in Monetary and Fiscal Policies. *Brookings Papers on Economic Activity*, 2, 139-215.
11. Romer C.D. (2000). Federal Reserve Information and the Behavior of Interest Rates. *American Economic Review*, 90 (3), 429-457.
12. Rotemberg J., Woodford M. (1999). Interest Rate Rules in an Estimated Sticky Price Model. W: *Monetary Policy Rules*. Ed. J.B. Taylor. University of Chicago Press, Chicago.
13. Sargent T., Wallace N. (1981). Some Unpleasant Monetarist Arithmetic. *Federal Reserve Bank of Minneapolis Quarterly Review*, 5, 1-17.
14. Szpunar P. (2000). *Polityka pieniężna. Cele i warunki skuteczności*. PWE, Warszawa.
15. Taylor J.B. (1993). Discretion versus Policy Rules in Practice. *Carnegie-Rochester Conference Series on Public Policy*, 39, 195-214.
16. Walsh C. (2001). Transparency in Monetary Policy. *FRBSF Economic Letter* 2001, 26.

17. Wojtyna A. (1996). Niezależność banku centralnego a teoretyczne i praktyczne aspekty koordynacji polityki pieniężnej i fiskalnej. *Bank i Kredyt*, 6.
18. Wojtyna A. (1998). *Szkice o niezależności banku centralnego*. PWN, Warszawa
19. Woroniecka I. (2004). Factors determining interest rate level in Poland. Estimation results for 1993-2002. W: *MODEST 2004: Integration, Trade, Innovation and Finance: From Continental to Local Perspectives*. Red. J.W. Owsinski. Polish Operational and Systems Research Society, Warszawa, 21-40.
20. Woroniecka I. (2006). Gra o politykę makroekonomiczną między bankiem centralnym a rządem. W: *Badania operacyjne i systemowe 2006. Analiza systemowa w globalnej gospodarce opartej na wiedzy: e-wyzwania*. Red. E. Urbaczyk, A. Straszak, J.W. Owsinski. Akademicka Oficyna Wydawnicza EXIT, Warszawa, 153-166.
21. Woroniecka I. (2007): Analiza priorytetów banku centralnego w polityce stóp procentowych. *Ekonomista*, 4, 559-580.
22. Woroniecka I. (2008): Pareto-optymalność rozwiązań w grze między bankiem centralnym a rządem. W: *Modelowanie preferencji a ryzyko '08*. Red. T. Trzaskalik. AE, Katowice, 127-142.
23. Woroniecka-Leciejewicz I. (2008). Dylemat więźnia i inne przypadki grze monetarno-fiskalnej. W: *Badania operacyjne i systemowe: decyzje, gospodarka, kapitał ludzki i jakość*. Red. J.W. Owsinski, Z. Nahorski, T. Szapiro. *Badania Systemowe*. IBS PAN, Warszawa, 64, 161-172.