



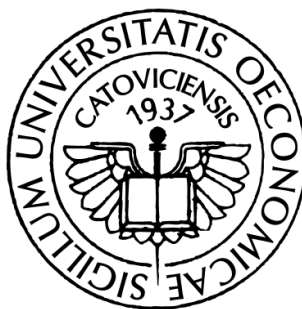
# **Modelowanie preferencji a ryzyko '21-'22**

## **Praca naukowa**



# **Modelowanie preferencji a ryzyko '21-'22**

**Redakcja naukowa**  
**Maciej NOWAK**  
**Tadeusz TRZASKALIK**



**Katowice 2021**

## **Komitety Redakcyjny**

Janina Harasim (przewodnicząca), Monika Ogrodnik (sekretarz),  
Małgorzata Pańkowska, Jacek Pietrucha, Irena Pyka, Anna Skórska,  
Maja Szymura-Tyc, Artur Świerczek, Tadeusz Trzaskalik, Ewa Ziemia

## **Recenzent**

Ewa Roszkowska

## **Redakcja i korekta językowa**

Karolina Koluch

## **Skład tekstu**

Daria Liszowska

## **Projekt okładki**

Janusz Gumulak

Ilustracja na okładce © wykres autorstwa Józefa Stawickiego

## **Druk i oprawa**

EXDRUK Spółka Cywilna Wojciech Żuchowski, Adam Filipiak  
ul. Rysia 6, 87-800 Włocławek  
e-mail: [biuroexdruk@gmail.com](mailto:biuroexdruk@gmail.com), [www.exdruk.com](http://www.exdruk.com)

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>

Egzemplarze elektroniczne publikacji mogą być również oferowane w komercyjnych kanałach dystrybucji.  
Udostępnienie ich w takich kanałach nie wpływa na otwarty charakter publikacji ani na prawa użytkowników  
wynikające z powyższej licencji

**ISBN 978-83-7875-779-5**  
**e-ISBN 978-83-7875-780-1**

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2021



WYDAWNICTWO UNIwersytetu Ekonomicznego w KATOWICACH

ul. 1 Maja 50, 40-287 Katowice, tel.: +48 32 257-76-33  
[www.wydawnictwo.ue.katowice.pl](http://www.wydawnictwo.ue.katowice.pl), e-mail: [wydawnictwo@ue.katowice.pl](mailto:wydawnictwo@ue.katowice.pl)  
Facebook: [@wydawnictwouekatowice](https://www.facebook.com/wydawnictwouekatowice)



## Spis treści

|                                                                |   |
|----------------------------------------------------------------|---|
| <b>Wstęp</b> ( <i>Maciej Nowak, Tadeusz Trzaskalik</i> ) ..... | 7 |
|----------------------------------------------------------------|---|

### Część I

#### ANALIZA PREFERENCJI I ANALIZA RYZYKA

|                                                                                                                                             |    |
|---------------------------------------------------------------------------------------------------------------------------------------------|----|
| 1. <b>Hybrydyzacja metod System Dynamics i Soft System Dynamics w celu wspomagania decyzji menedżerskich</b> ( <i>Dariusz Banaś</i> ) ..... | 13 |
| 2. <b>Zastosowanie cen online i modeli autoregresyjnych w krótkoterminowych prognozach cen żywności</b> ( <i>Jarosław Janecki</i> ) .....   | 43 |
| 3. <b>Teoria ujawnionych preferencji i wybory międzyokresowe w warunkach ryzyka</b> ( <i>Józef Stawicki</i> ) .....                         | 52 |

### Część II

#### RYZYKO W ZARZĄDZANIU PROJEKTAMI

|                                                                                                                                                  |    |
|--------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 1. <b>Zarządzanie ryzykiem projektu w małych instytucjach kultury z wykorzystaniem Scruma</b> ( <i>Dorota Kuchta, Alicja Krawczyńska</i> ) ..... | 67 |
| 2. <b>Sprężenie zwrotne w adaptacyjnym zarządzaniu ryzykiem projektu</b> ( <i>Dariusz Meiser</i> ) .....                                         | 83 |
| 3. <b>Propozycja globalnego wdrożenia systemu ERP z wykorzystaniem podejścia obiektowego</b> ( <i>Łukasz Tync</i> ) .....                        | 95 |

### Część III

#### UCZENIE MASZYNOWE

|                                                                                                                                                                  |     |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 1. <b>Zastosowanie uczenia maszynowego w predykcji rezultatów meczów piłki nożnej</b> ( <i>Szymon Głowania, Jan Kozak</i> ) .....                                | 119 |
| 2. <b>Rozwój wyjaśniania modeli uczenia maszynowego a zaufanie do sztucznej inteligencji</b> ( <i>Szymon Głowania, Bogna Zacny, Grzegorz Dziczkowski</i> ) ..... | 134 |





## Wstęp

Analiza i modelowanie preferencji, jak również modelowanie ryzyka to problematyka z zakresu badań operacyjnych znajdująca się na styku ekonomii, zarządzania, matematyki stosowanej i informatyki. Jest ona przedmiotem zainteresowania Autorów niniejszej monografii, która została podzielona na trzy części. W kolejnych jej rozdziałach przedstawiono zagadnienia z zakresu analizy preferencji i analizy ryzyka, zarządzania ryzykiem przy realizacji projektów oraz wybranych problemów uczenia maszynowego. Poniżej zaprezentowano główne problemy podjęte w kolejnych rozdziałach monografii.

### I. Analiza preferencji i analiza ryzyka

Jedna z międzynarodowych firm z branży IT o zasięgu transnarodowym, w wyniku strategicznej orientacji na długofalowe dostarczanie innowacyjnych usług telekomunikacyjnych, zamierza wprowadzić na rynek pakiety usług sieci operatorskich typu 5G2 oraz przygotować się pod 6G3. Problemem badawczym rozpatrywanym w tak sformowanym zadaniu biznesowym jest dynamiczne oszacowanie możliwości pozyskania klientów w obszarze akwizycji, późniejszej sprzedaży usług oraz zwymiarowanie organizacji dostarczania usług pod względem operacyjnym w relacji do spodziewanej absorpcji tychże usług przez operatorów sieci telekomunikacyjnych w horyzoncie 60 miesięcy. Przedstawione w rozdziale *Hybrydyzacja metod System Dynamics i Soft System Dynamics w celu wspomagania decyzji menedżerskich* (Dariusz Banaś) podejście hybrydyzacji otwiera dodatkowe możliwości metodyczne polegające na syntezie modeli dynamiki systemowej SD oraz SSD WINGS w środowisku BSC (Zbilansowanej Karty Wyników).

Rosnąca popularność zawierania transakcji online, wzrost wartości transakcji przeprowadzanych online oraz możliwość porównywania cen, jakie oferuje internet, stwarzają nowe możliwości monitorowania zmian cen. Niepewność, z jaką mamy do czynienia w przypadku krótkoterminowych prognoz wskaźnika inflacji (prognozy w ramach konsensów rynkowych a dane publikowane przez GUS), nie jest czynnikiem pożądanym. Ceny online stwarzają możliwość uży-

skania szybszej i pełniejszej wiedzy o procesach cenowych. Problemy te poruszono w rozdziale *Zastosowanie cen online i modeli autoregresyjnych w krótkoterminowych prognozach cen żywności* (Jarosław Janecki).

W rozdziale *Teoria ujawnionych preferencji i wybory międzyokresowe w warunkach ryzyka* (Józef Stawicki) podjęto problemy z obszaru teorii ujawnionych preferencji i międzyokresowego podejmowania decyzji. Są to teorie dobrze rozpoznane. Do ujawnionych preferencji potrzeba dynamicznego podejścia. Słaby aksjomat ujawnionych preferencji Afriata nie wystarcza do określenia doskonałej zgodności między obserwowanymi wyborami a nieobserwowalnymi ujawnionymi preferencjami. Niezbędne i konieczne jest określenie warunków koniecznych i wystarczających dla tej zgodności. Warunki niedoskonałej zgodności są analizowane w niektórych pracach. Ustalenia zawarte w dotychczasowej literaturze stanowią niezbędny pierwszy krok w kierunku głębszych badań dla docenienia związku między obserwowanymi dynamicznymi wyborami a dynamicznie ujawnionymi preferencjami. Dynamiczne wybory to wybory międzyokresowe (decyzje), w których rozłożenie w czasie kosztów i korzyści zmienia się w czasie. Są one zarówno powszechne, jak i ważne. Szczególnie pytanie stawiane nie tylko dziś to na przykład, ile odkładać na emeryturę czy jak inwestować. Pytania te są istotnymi pytaniami, a decyzje w tym obszarze mają mocne podstawy w teorii wyborów międzyokresowych.

## II. Ryzyko w zarządzaniu projektami

W rozdziale *Zarządzanie ryzykiem projektu w małych instytucjach kultury z wykorzystaniem Scruma* (Dorota Kuchta, Alicja Krawczyńska) omówiono charakterystykę działalności instytucji kultury na szczeblu samorządowym w kontekście specyfiki podejmowanych projektów, przedstawiono podstawowe pojęcia związane z zarządzaniem ryzykiem oraz zaprezentowano pojęcie Scruma jako pewnego rodzaju zbioru zasad dotyczącego prowadzenia projektów oraz zakres procesów zarządzania ryzykiem. Ostatnia część rozdziału to opis autorskiego modelu zarządzania ryzykiem dostosowanego do potrzeb instytucji kultury.

W rozdziale *Sprzężenie zwrotne w adaptacyjnym zarządzaniu ryzykiem projektu* (Dariusz Meiser) omówiono koncepcje „multisprzężenia zwrotnego” oraz „samosprzężenia zwrotnego” w metodzie adaptacyjnego zarządzania ryzykiem projektu wytwórczego. „Multisprzężenie zwrotne” dotyczy sytuacji, gdy w schemacie blokowym proponowanej metody sprzężenia zwrotne występują od

każdego bloku do każdego innego bloku poprzedzającego; nie tylko do bloku poprzedniego, ale również do każdego wcześniejszego bloku. Z kolei „samosprężenie zwrotne” polega na wprowadzeniu sprzężenia zwrotnego z danego bloku do tego samego bloku, w sposób iteracyjny, dopóki nie zostanie spełniony jeden z warunków: flaga „ON”, pozwalająca przejść do następnego bloku, wtedy gdy dany warunek zostanie spełniony lub „timeout” – upływie czas lub liczba iteracji, jakie zostały przewidziane dla danego warunku. W rozdziale przedstawiono te dwie koncepcje sprzężeń zwrotnych w ramach prac autora nad metodą adaptacyjnego zarządzania ryzykiem wytwórczym w branży urządzeń elektronicznych.

W rozdziale *Propozycja globalnego wdrożenia systemu ERP z wykorzystaniem podejścia obiektowego* (Łukasz Tync) rozpatrywany jest program składający się z wielu podobnych projektów informatycznych oraz wypływającej z nich transformacji procesów biznesowych realizowanych w zróżnicowanym środowisku organizacyjnym. Celem rozdziału jest budowa modelu otoczenia programu globalnego wdrożenia systemu ERP ze szczególnym uwzględnieniem interfejsów komunikacyjnych niezbędnych do prawidłowego osadzenia tego programu w organizacji. Przykładowy projekt jest fuzją wieloletnich doświadczeń, które autor zebrał jako menedżer i konsultant, uczestnicząc w podobnych programach.

### III. Uczenie maszynowe

Predykcja wyników meczów piłki nożnej jest znanym problemem, brakuje jednak rozwiązań związanych z kompleksowym podejściem do odkrywania wiedzy z danych. Jest to w szczególności trudne w przypadku mniej popularnych lig piłkarskich. W rozdziale *Zastosowanie uczenia maszynowego w predykcji rezultatów meczów piłki nożnej* (Szymon Głowania, Jan Kozak) zaproponowano autorskie podejście do całego procesu odkrywania wiedzy z danych pozwalającego na przewidywanie rezultatów spotkań polskich lig piłkarskich. Celem rozdziału jest wybranie algorytmów uczenia maszynowego pozwalających na poprawę predykcji wyników meczów piłkarskich w stosunku do przewidywań bukmacherów. Zaproponowane podejście przetestowano na rzeczywistych danych – wynikach spotkań piłkarskich polskiej PKO Ekstraklasy.

Aplikacje oparte na sztucznej inteligencji, w szczególności na technikach i algorytmach uczenia maszynowego, stają się coraz bardziej popularne i dotyczą coraz szerszych obszarów ludzkiej działalności. Technologie te znajdują zastosowanie w coraz ważniejszych i wrażliwszych dziedzinach naszego życia.

Aby móc w pełni wykorzystać systemy oparte na sztucznej inteligencji, konieczne staje się określenie naszego zaufania do wyników poszczególnych algorytmów. Celem rozdziału ***Rozwój wyjaśniania modeli uczenia maszynowego a zaufanie do sztucznej inteligencji*** (Szymon Głowania, Bogna Zacny, Grzegorz Dziczkowski) jest zwrócenie uwagi na bardzo istotny aspekt zaufania do wyników inteligentnych algorytmów. Autorzy opisują ewolucję terminu „zaufanie”, następnie przedstawiają otrzymaną wiedzę z algorytmów uczenia maszynowego, finalnie skupiając się na interpretowalności i wytłumaczalności algorytmów koniecznych do otrzymania godnej zaufania sztucznej inteligencji.

Maciej Nowak  
Tadeusz Trzaskalik

## Część I



## **ANALIZA PREFERENCJI I ANALIZA RYZYKA**



# Hybrydyzacja metod System Dynamics i Soft System Dynamics w celu wspomagania decyzji menedżerskich

*Dariusz Banaś<sup>1</sup>*

## Wprowadzenie

Jedna z międzynarodowych firm z branży IT<sup>2</sup> o zasięgu transnarodowym, w wyniku strategicznej orientacji na długofalowe dostarczanie innowacyjnych usług telekomunikacyjnych, zamierza wprowadzić na rynek pakiety usług sieci operatorskich (elementy portfolio serwisowego – EP) typu 5G<sup>3</sup> oraz przygotować się pod 6G<sup>4</sup>. Problemem badawczym rozpatrywanym w tak sformułowanym zadaniu biznesowym jest dynamiczne oszacowanie (z użyciem adekwatnych metodyk badawczych do dynamicznej analizy operacyjnych problemów decyzyjnych) możliwości pozyskania klientów w obszarze akwizycji, późniejszej sprzedaży usług oraz zwymiarowanie organizacji dostarczania usług pod względem operacyjnym w relacji do spodziewanej absorpcji tychże usług przez operatorów sieci telekomunikacyjnych w horyzoncie 60 miesięcy. Pierwszorzędnym celem biznesowym i celem badania dla tak sformułowanego problemu badawczego jest osiągnięcie dochodowości na założonym poziomie oraz zapewnienie akceptowalnego kompromisu jakości oferowanych usług, wpływającego na koszty własne. Metodycznie, z uwagi na spektrum zagadnień z obszarów finansowych i niefinansowych, danych jakościowych i ilościowych, a także natywnej dynamiki procesów występujących w obszarach problemu badawczego i celu badania, zastosowano hybrydowy model dynamiki systemowej (System Dyna-

---

<sup>1</sup> Uniwersytet Ekonomiczny w Katowicach, e-mail: [dariusz.banas@edu.uekat.pl](mailto:dariusz.banas@edu.uekat.pl).

<sup>2</sup> IT (ang. Information Technology) to technologia informacyjna będąca podzbiorem technologii informacyjno-komunikacyjnych (ICT) [Axelos, 2011]. Termin IT jest używany zwykle w kontekście informatycznej działalności biznesowej, obejmując procesy (planowanie, eksploatację, kontrolę, zarządzanie), sprzęt, oprogramowanie (infrastrukturę), usługi.

<sup>3</sup> W ramach technologii mobilnej piątej generacji.

<sup>4</sup> W ramach technologii mobilnej szóstej generacji.

mics – SD) oraz miękkiej dynamiki systemowej (Soft System Dynamics – SSD WINGS<sup>5</sup>) do zaprojektowania możliwych scenariuszy rozwojowych w okresie 5-letnim w celu wsparcia decyzji menedżerskich. Przedstawione podejście hybrydyzacji otwiera dodatkowe możliwości metodyczne w badaniach operacyjnych polegające na syntezie modeli dynamiki systemowej SD oraz SSD WINGS w środowisku BSC (Ballanced Scorecard/ Zbilansowanej Karty Wyników).

W analizie powyższych zagadnień (także jako nawiązaniu do ujęcia teoretycznego – modelowego przedstawionego w podrozdziale 1.1), należy wziąć pod uwagę dodatkowe czynniki, takie jak środowisko, w którym znajduje się obecnie uogólnione przedsiębiorstwo<sup>6</sup>, a także kadrę menedżerską oraz pracowników firmy<sup>7</sup>. Przedsiębiorstwo jako obiekt gospodarczy jest zbiorem różnych struktur (ludzkich, zasobowych, finansowych, informacyjnych itd.) oraz procesów (główne: klienta, produktu, dostawy, zarządzania i wsparcia), co w praktyce operacyjnej przekłada się (w najlepszym przypadku – w rzeczywistości firm<sup>8</sup>) na monitorowanie stanu przedsiębiorstwa w perspektywach Zbilansowanej Karty Wyników (BSC – Balanced Scorecard) [Kaplan, Norton, 1992] – bez analizy dynamicznych współzależności występujących między tymi perspektywami<sup>9</sup>.

Ze względu na wyżej wymienioną „złożoność” przedsiębiorstwa skuteczne metody analizy i oceny jego działania w celu podejmowania decyzji powinny próbować uwzględniać całościowo rzeczywisty charakter obiektu ekonomicznego – przedsiębiorstwa, czyli organizację oraz jej środowisko operacyjne, z uwzględnieniem ich cech charakterystycznych, które nie są: liniowymi, statycznymi bądź będącymi w stanie równowagi obiektami (gdzie rzeczywiste zachodzące procesy nie są również pozbawione opóźnień). Próba stworzenia chociażby podstawowego, lecz całościowego modelu organizacji wymaga zatem zrozumienia jego „eigenvalues” w obszarach procesów biznesowych, interakcji wewnętrznych i zewnętrznych, wpływu otoczenia itp. poprzez różne pryzmaty funkcjonalne zaangażowanych w nich ludzi oraz dostępność wiarygodnych da-

<sup>5</sup> Tu: SSD użyte w kontekście metody WINGS – patrz szczegóły w rozdziale 1.1.

<sup>6</sup> Stające się coraz bardziej złożonym, gdzie informacje są nieodzownym elementem jego otoczenia wewnętrznego i zewnętrznego, rozumianego jako system.

<sup>7</sup> Muszą oceniać sytuację nie tylko z punktu widzenia kondycji firmy per se (gdzie informacje są gromadzone w wielu obszarach), ale i również sine qua non przez pryzmat swoich osobowych możliwości.

<sup>8</sup> Według własnych obserwacji, a także na bazie współpracy z firmami konsultingowymi, np. Gartner, Bain.

<sup>9</sup> Podejścia z dynamicznym/nieliniowym opisem wzajemnych relacji modelu/obiektu i wykorzystaniem również BSC nie są praktykowane (przynajmniej w rozpatrywanym przypadku tematu podjętego w niniejszym rozdziale i firmie jako takiej).

nych z różnych obszarów działania firmy. Próba powyższa – oparta na zintegrowaniu modeli myślowych zaangażowanej kadry (ich rozumienia spraw, interpretacji, osądów – wiedzy eksperckiej) – rozumiana jako ujęcie jakościowe, wraz z ogólnym myśleniem i modelowaniem systemowym i danymi „twardymi” – jako ujęcie ilościowe, stanowi podstawę do zbudowania modelu dynamicznego i zwymiarowania organizacji dostarczania usług pod względem operacyjnym w relacji do liczby klientów zainteresowanych zakupem usług 5G (i w przyszłości 6G) oraz wspomagania decyzji menedżerskich przy ocenie wpływu wprowadzenia nowych usług.

Według analiz oraz obecnego stanu wiedzy (podrozdział 1.1) powyższe może być osiągnięte w wyniku jednoczesnego użycia (hybrydyzacji) metodyk „miękkich” (dane jakościowe) i „twardych” (dane ilościowe) modelowania systemowego. Odzwierciedla to zastosowanie miękkich badań operacyjnych (soft system methodology) reprezentowanych przez soft system dynamics w połączeniu z tzw. metodami „twardymi” w ujęciu perspektyw BSC. Nie bez znaczenia (w ujęciu praktycznym) jest także wymaganie sterowania takim modelem, które powinno być najpierw wieloaspektowym monitorowaniem (wraz z wyznaczeniem tzw. benchmarków operacyjnych), a następnie zawierać wypracowanie możliwych rozwiązań i ich zastosowanie w celu osiągnięcia zamierzonych (wybranych) rezultatów.

## 1.1. Hybrydyzacja – część teoretyczna

W kontekście hybrydyzacji krótkie podsumowanie, czym jest, a czym nie SD i SSD, zamieszczono w tabeli 1, a poszerzające (uzupełniające) informacje teoretyczne o tematyce SD i SSD zawarto w podrozdziale 1.3.

Na podstawie studium literatury w temacie hybrydyzacji dynamiki systemowej (SD i SSD) otrzymano następujące wnioski (1-10):

1. Według przyjętego nazewnictwa [Mendoza, Prabhu, 2006] SSD można zaliczyć do grupy SSM (Soft System Methodology [Checkland, 2000]) – choć nie bezpośrednio, gdyż np. metoda FCM (Fuzzy Cognitive Mapping) [Kosko, 1986] jest specyficzną wersją SSM [Hanafizadeh, Aliehyaei, 2011].
2. WINGS (SSD) [Adamus-Matuszyńska, Michnik, Polok, 2019; Banaś, Michnik, 2019] jako jej rozszerzona dynamiczna forma w porównaniu do oryginalnej metody [Michnik, 2013] klasyfikuje się wprost jako SSD.

3. Próby łączenia SSM oraz SD występujące w literaturze [Lane, Oliva, 1998] prowadzą do powstania nowej klasy „metodologii” nazwanej Soft System Dynamics Methodology (SSDM) [Rodríguez-Ulloa, Paucar-Caceres, 2005; Rodríguez-Ulloa, Montbrun, Martínez-Vicente, 2011] o skomplikowanej ramie metodycznej<sup>10</sup>, jednakże SSM w tym wykonaniu to nie np. metoda FCM, a (ogólniej) soft OR.
4. „W literaturze multimetodologicznej istnieją kontrowersje dotyczące paradygmatycznego i konceptualnego statusu mieszania metod” [Tashakkori, Teddlie, 2010] – lecz w obecnym stanie rzeczy jest to raczej dyskusja formalna – niewskazująca na jakieś konkretne rozstrzygnięcie (gdyż raczej go być nie może).
5. W ujęciu badań operacyjnych brakuje metod w pełni dynamicznych, które jednocześnie uwzględniają modelowanie ilościowe (w ujęciu SD) i jakościowe (SSD) [Kosko, 1986; Mendoza, Prabhu, 2006; Michnik, 2013], łącząc w sobie mocne strony obydwu metod oraz kompensując wady każdej z nich [Coyle, 2000; Forrester, 2013].
6. Mimo iż występuje ten sam cel, zarówno w modelach ilościowych (SD), jak i jakościowych (SSM), opisanych jest niewiele prób scalenia metod w ogóle, według poszukiwań własnych i dostępnych analiz (np. na podstawie studium literatury [Zolfagharian, Romme, Walrave, 2018]) na temat łączenia SD z inną metodą [Zolfagharian, Romme, Walrave, 2018] w celu podejmowania decyzji.
7. Zidentyfikowane konkretne (w celu podejmowania decyzji) przykłady scalenia metod to np. połączona procedura SD oraz PROMETHEE [Zolfagharian, Romme, Walrave, 2018]<sup>11</sup>, jak również inne podejścia dynamiczne [Kasperska, 2005; Rabelo i in., 2005], choć prace w kierunku hybrydyzacji trwają już od dawna (także) z konkretnymi przypadkami różnych implementacji [Stermann, 2000, s. 96]: „Typowym przykładem jest System Oceny Niezależności Projektów – hybrydowy model oparty na programowaniu liniowym, ekonometrii oraz analizie wejść/wyjść”.

Każdy przypadek jest jednak indywidualny w zależności od rozpatrywanego problemu (nie ma jednolitego wzorca hybrydyzacji).

---

<sup>10</sup> „Iteracyjny charakter struktury SD może nawet sugerować decyzję o rezygnacji z modelowania SD, zwłaszcza gdy inne metody (połączone lub samodzielne) wydają się bardziej odpowiednie i obiecujące” [Rodríguez-Ulloa, Montbrun, Martínez-Vicente, 2011].

<sup>11</sup> SD: Symulacja scenariuszy, opcji strategicznych i test długoterminowy; MCDA: Ranking i wybór za pomocą procedury PROMETHEE.

8. Chociaż istnieje ogólna zgoda co do znaczenia danych jakościowych podczas opracowywania modelu dynamiki systemu, nie ma również jasnego/ogólnego opisu, w jaki sposób i kiedy można ich używać [Luna-Reyes, Andersen, 2003], gdyż zależy to od konkretnego przypadku użycia<sup>12</sup>, stąd temat wydaje się raczej problematyczny w kierunku poszukiwań unitaryzacji.
9. W kontekście tematyki hybrydyzacji przedstawionej w rozdziale metody dynamiki systemowej i miękkiej dynamiki systemowej (rozpatrywane oddzielnie) rozwiązują postawione przed nimi problemy selektywnie. Hybrydyzacja SD i SSD wnosi nową jakość metodyczną do wspierania podejmowania kluczowych decyzji biznesowych w środowisku przedsiębiorstw (oraz modelowania np. scenariuszy rozwojowych), kreując „zsyntezowaną” metodę odpowiadającą cechom rozpatrywanego problemu – natywnym cechom przedsiębiorstwa, jego otoczenia i postępujących zmian transformacyjnych.
10. Hybrydyzacja SD i SSD wnosi także nową jakość metodyczną w badaniach operacyjnych, kreując bardziej adekwatną metodę odpowiadającą rzeczywistym cechom otaczającego świata – nie istnieje obecnie hybrydowe ujęcie (ilościowo-jakościowe), które łączy w sobie mocne strony obydwu metodyk/kompensuje wady każdej z nich.

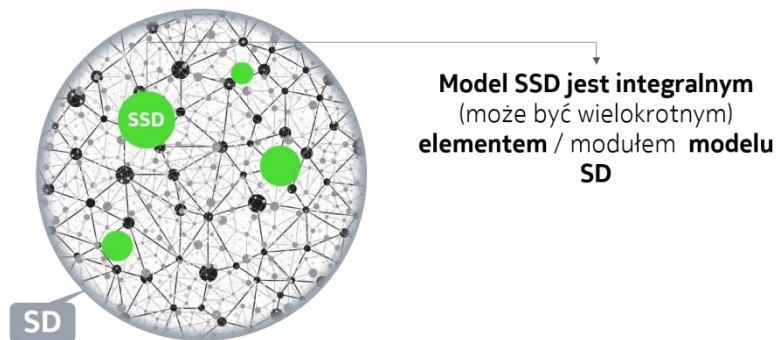
**Tabela 1.** Ogólne podsumowanie cech twardej i miękkiej dynamiki systemowej

| System Dynamics (SD)                                                                                                                                                                                                        | Soft System Dynamics (SSD)                                                                                                                                                     |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| SD to „twarda” metoda i technika modelowania matematycznego, próba zrozumienia zachowania złożonych nieliniowych systemów dynamicznych, gdzie zachowanie całości nie może być wyjaśnione na podstawie zachowania się części | SSD to metoda „mięka”, intuicja i doświadczenie, czyli korzystanie z „niejawnych” modeli (mental models). Metoda ta wykorzystuje głównie dane jakościowe (ale także ilościowe) |
| WADY: każda relacja w modelu musi być określona ilościowo – transpozycja jakościowa jest obciążona arbitralnością                                                                                                           | WADY: ograniczone możliwości symulacyjne i predykcyjne spowodowane większą niepewnością (wynikającą z niepewności źródeł danych). Odwzorowania dynamiczne ograniczone          |
| ZALETY: możliwości symulacyjne, predykcyjne / modelowanie scenariuszowe, wizualizacja ilościowa                                                                                                                             | ZALETY: możliwość transpozycji wiedzy niejawnej (ekspertkiej) oraz danych ilościowych w logiczną strukturę formalną                                                            |

<sup>12</sup> Na przykład gdy model wymaga użycia takich zmiennych miękkich, jak: „zadowolenie klienta”, „jakość produktu”, „presja na obniżenie ceny”, „zaangażowanie”, „postrzegana produktywność” czy też „bieg/ jakość procesów wewnętrznych” w powiązaniu z danymi ilościowymi.

## Hybrydyzacja – budowanie modelu

Zastosowanie SD w kontekście hybrydyzacji i przedsiębiorstwa to „przeniesienie” procesów przedsiębiorstwa w ramy układu monitorowania i regulacji, gdzie są obserwowane parametry w obszarach perspektyw Zbilansowanej Karty Wyników (BSC). W tak ustanowionym środowisku następuje analiza danych i są dokonywane symulacje zachowania przedsiębiorstwa w okresie długoterminowym (opisane w podrozdziale 1.2 – część empiryczna): okres 5 lat od  $t = 0$  do  $t = 60$  (w ujęciu miesięcznym), oraz wypracowywane korekty przyjętych działań operacyjnych (lub też ewentualnie ustanawiane nowe wytyczne w zależności od zachowania ekosystemu przedsiębiorstwa i bieżących warunków zewnętrznych). Zastosowanie BSC przez pryzmat dynamiki systemowej [Akkermans, Oorschot, 2002; Barnabè, 2011; Linard, Yoon, 2000; Nielsen, Nielsen, 2015], oprócz działań operacyjnych, to również możliwość dynamicznego dopasowywania strategii przedsiębiorstwa do zmieniającego się otoczenia, gdzie w sposób ciągły może być sprawdzana realizacja celów (działania takie pozwalają również wychwycić cele, których nie można zrealizować). System Dynamics (SD) to zatem „droga” do zrozumienia dynamicznego zachowania złożonego systemu, którym jest przedsiębiorstwo i którego zachowanie wynika z jego wewnętrznej struktury. „Struktura ta składa się z pętli sprzężenia zwrotnego, zapasów i przepływów oraz nieliniowości powstałych w wyniku interakcji fizycznej i instytucjonalnej struktury systemu z procesami decyzyjnymi działającymi w nim podmiotów” [Stermen, 2000]. Zastosowanie SSD w kontekście hybrydyzacji to „umieszczenie” SSD w strukturze SD<sup>13</sup> (rys. 1<sup>14</sup>).



**Rys. 1.** Model SSD jako część modelu SD

<sup>13</sup> Może też wystąpić sytuacja odwrotna, tj. SD jest częścią SSD, jednakże w tej konkretnej realizacji hybrydowego modelu wspomagania decyzji menedżerskich łączącego dynamikę systemową i miękką dynamikę systemową z uwagi na przyjęte cele operacyjne i adekwatną do tego budowę modelu SD/SSD nie zachodzi taka realizacja.

<sup>14</sup> Uproszczony schemat ilustrujący realizację modelu hybrydowego SD/SSD.

1. **Hybrydowy model wspomaganie decyzji menedżerskich** łączący dynamikę systemową (SD) i miękką dynamikę systemową (SSD) – **to hybrydyzacja metody SD Forrestera** [Forrester, 2013] oraz **SSD WINGS Michnika** [Banaś, Michnik, 2019] z modelowaniem (wypracowaniem) możliwych scenariuszy/ścieżek rozwojowych w celu adaptacyjnego wspomaganie decyzji menedżerskich w firmie w środowisku BSC, tworząc przez to **uniwersalne podejście metodyczne do syntezy danych jakościowych i ilościowych**.

Hybrydowy model SD/SSD przedstawiony w rozdziale stanowi konkretny przypadek zastosowania w kontekście użyteczności do pakietu rozwiązań serwisowych – usług EP (Elementów Portfolio), składającego się z 21 elementów (EP1, EP2, EP3, ... EP21). Celem takiego podejścia i wykorzystania hybrydowego modelu SD/SSD jest dostosowanie operacyjne organizacji firmy (z użyciem danych jakościowych i ilościowych, które odzwierciedlają specyfikę procesów dla usług EP) w relacji do spodziewanej absorpcji tychże usług przez operatorów sieci telekomunikacyjnych na poziomie jakości będącym kompromisem pomiędzy: maksymalną możliwą jakością (przy znacznie zwiększonym koszcie firmy) a aprobowalną przez klientów – na podstawie pomiaru stopnia ryzyka (ich niezadowolone – rezygnacji z aktywnego używania usług, a co za tym idzie potencjalnie mniejszym zainteresowaniem tego typu usługami w przyszłości), z pierwszorzędnym celem osiągnięcia dochodowości firmy na założonym poziomie ze sprzedaży EP. Rozwiązania serwisowe EP trafią do kanałów sprzedażowych oraz implementacji, gdzie (w połączonym modelu SD-SSD) w procesie jakościowym (to proces sprzedaży) poprzez użycie metody WINGS jest obliczana miara prawdopodobieństwa sukcesu procesu ofertowego/sprzedaży EP – oznaczana w modelu jako wskaźnik (nominalny) WINGS EP (jako odniesienie: jest on również zasygnalizowany w podrozdziale 1.3.2). Liczba sprzedanych elementów portfolio, które trafią do implementacji u klienta, ostatecznie zależy od: wskaźnika nominalnego WINGS EP, liczby potencjalnych klientów firmy przystępujących do procesu sprzedaży usług (kupna w rozumieniu klienta) oraz od opóźnienia (losowego z przedziału 6-9 miesięcy – równanie 7), które wnosi proces sprzedażowy (wraz z procesem przygotowania do dostarczania usług klientowi przez organizację Global i Local Service Delivery<sup>15</sup>).

**Dla powyższego zostanie przetestowany wpływ wprowadzenia pakietu serwisowego EP na zwymiarowanie operacyjne firmy w celu dostarczania usług wraz z aspektami finansowymi i wypracowania niezbęd-**

---

<sup>15</sup> Globalne i lokalne zespoły realizacji usług EP 5G.

nych decyzji w dziedzinie zarządzania operacjami firmy (również strategicznych w kontekście innych przyszłych rozwiązań, np. dla sieci 6 [Vetter, Viswanathan, 2021]).

2. **Hybrydyzacja SD i SSD zachodzi na poziomie obliczeń maszynowych** poprzez: synchroniczne „wstrzyknięcie” 21 wektorów „Wskaźnika sukcesu ofertowego – sprzedaży EP (WINGS EP)” wyznaczonych dla każdego EP (każdy wektor zawiera 61 – od  $t = 0$  do  $t = 60$  – elementów danych) z jakościowego modelu WINGS (rys. 3<sup>16</sup>) do 21 modułów ilościowego modelu SD (na jeden moduł przypada jeden wektor; moduły są oznaczone jako „prostokąty” z „zaokrąglonymi rogami”). Połączony model SD/SSD przedstawiono na rys. 2<sup>17</sup>.

„Wstrzykiwanie” danych odbywa się do modułów **na każdy cykl obliczeniowy w danej chwili czasu  $t$  z krokiem  $dt$** . Krok  $dt$  o wartości wynoszącej 1 (czyli 1 miesiąc) został dobrany analitycznie<sup>18</sup> do specyfiki układu SD/SSD na podstawie twierdzenia WKS<sup>19</sup>. Użyta metoda całkowania numerycznego to RK4 (Runge-Kutta 4 z uwagi na jej większą dokładność w modelu SD/SSD niż np. metody Eulera).

W celu wygenerowania wyżej wymienionych 21 oddzielnych wektorów<sup>20</sup> należy najpierw obliczyć wielkości nominalne WINGS EP (według sieci zależności przedstawionej na rys. 3) oddzielnie dla każdego wariantu od Idealnego do Antyidealnego, gdyż przyjęto, że przeprowadzone zostanie także modelowanie scenariuszowe w zależności od użytego wariantu (Antyidealny, Słaby, Średni, Dobry, Idealny). Wariant Antyidealny reprezentuje naj-

<sup>16</sup> Na rysunku tym wartości konceptów i interakcji zostały określone na podstawie wiedzy eksperckiej (dane jakościowe). Wskaźnik sukcesu ofertowego/sprzedaży EP (WINGS EP) jest wyznaczany dla każdego EP (Elementu Portfolio) oddzielnie w celu zasilenia danymi („wstrzyknięcie”) części SD modelu SD/SSD.

<sup>17</sup> Jest to szczegółowy schemat hybrydowego modelu SD/SSD w perspektywach BSC.

<sup>18</sup> Wraz z eksperymentalnym sprawdzeniem wartości stopy błędu generowanej przy innych wartościach  $dt$ , np. dla  $dt$  wynoszącego 1 tydzień, 3 dni i 1 dzień w wyniku analizy częstotliwości występujących w pętlach modelu SD/SSD. Maksymalna średnia wartość błędu dla  $dt = 1$  wynosiła  $\approx 5\%$  w ogóle dla przetestowanych różnych grup parametrów (średnio dla magazynów, przepływów i konwerterów).

<sup>19</sup> Twierdzenie o próbkowaniu WKS (Whittaker, Kotelnikov, Shannon); praktyczne aspekty dyskretyzacji [Gensun, 1996; Luke, 1999; Smith, 2007].

<sup>20</sup> Do wygenerowania 61 elementów wektora danych został wykorzystany „czynnik losowy”, który polega na użyciu równomiernego rozkładu prawdopodobieństwa dla zmiennych ciągłych (i dyskretnych). Rozkład ten został nałożony na wektor danych wyjściowych z WINGS (opisujący 60 miesięcy + chwila „0”) w zakresie od 0.9 do 1.1 wielkości nominalnej obliczonej z WINGS (czyli w zakresie 20%). Ma to urzeczywistnić charakter procesu sprzedaży w związku z założonym 20-procentowym czynnikiem losowym dla każdego wariantu/scenariusza realizacyjnego (od wariantu Antyidealnego do Idealnego). Kiedy wartość wylosowana przekraczała 1, była zastępowana liczbą 1 jako maksymalnym możliwym poziomem sukcesu.

mniejszą skuteczność sprzedaży, a zatem najmniejszą wartość „Wskaźnika sukcesu ofertowego – sprzedaży EP (WINGS EP)” – wariant Idealny przeciwnie. W modelu SD/SSD przekłada się to stosownie na najmniejszą liczbę sprzedanych EP (które trafią później do procesu dostarczania usług) oraz największą – dla wariantu Idealnego (z wartościami pośrednimi dla reszty wariantów).

Obliczenie wielkości nominalnych WINGS EP jest realizowane z sieci WINGS w wyniku interakcji pomiędzy conceptami<sup>21</sup> – zgodnie z ogólnym algorytmem obliczeniowym WINGS [Adamus-Matuszyńska, Michnik, Polok, 2019] (równania (13-16) w podrozdziale 1.3).

## 1.2. Hybrydyzacja – część empiryczna

### 1.2.1. Ilustracje realizacyjne

- Wszystkie wykresy przedstawione w niniejszym rozdziale zostały wykonane na podstawie symulacji hybrydowego modelu SD/SSD przedstawionego na rys. 2 [www 1].
- Schematyczny układ hybrydowy w perspektywach BSC (Klienta, Procesów Wewnętrznych, Wiedzy i Rozwoju, Finansowej) zamieszczono na rys. 1<sup>22</sup>, a szczegółowy na rys. 2<sup>23</sup> (gdzie detale części SSD są dostępne na rys. 3<sup>24</sup>).
- Realizacje układowe SD/SSD (wygenerowane jako pliki „xml version”/ „xmile version” oraz „isdb”<sup>25</sup>) dla każdego wariantu – Antyidealny, Słaby, Średni (oraz Średni bez fluktuacji związanych z „czynnikami losowymi”), Dobry, Idealny – wraz z odrębnymi „modelami danych” są dostępne pod linkiem [www 2].

<sup>21</sup> Dla conceptów (składników) systemu WINGS i wzajemnych interakcji użyto skali analogicznej do skali Likerta, jednakże w przeciwieństwie do oryginalnej skali (Likerta), która ma charakter interwałowy, ta użyta w WINGS jest skalą ilorazową (rozszerzoną również na wartości ujemne –5 do –1). Zastosowanie skali ilorazowej wynika z potrzeby przekształceń algebraicznych wykonywanych na ocenach eksperckich w zastosowanej metodzie SSD WINGS (oceny eksperckie różnią się w zależności od wariantu).

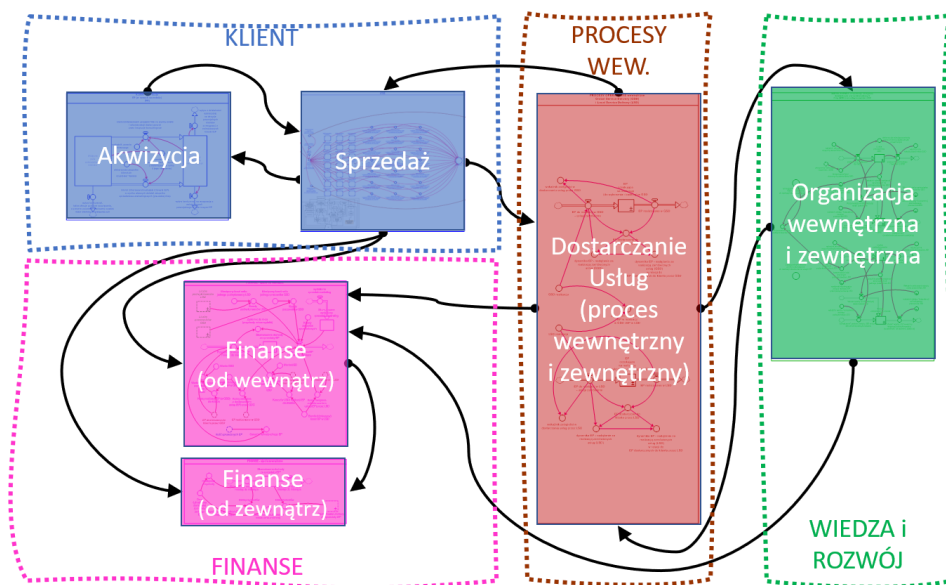
<sup>22</sup> Uproszczony schemat ilustrujący osadzenie modelu hybrydowego SD/SSD w perspektywach BSC.

<sup>23</sup> Rysunek ten odsyła za pośrednictwem linku <https> do materiału źródłowego w celu „zwiększenia czytelności” rysunku/ lepszej identyfikacji jego szczegółów.

<sup>24</sup> Lub pod linkiem [www 6].

<sup>25</sup> Plik „isdb” jest bazą danych iThink. iThink to narzędzie do modelowania dynamiki systemu opracowane przez isee systems.

- Wizualizacje dynamiczne całego modelu SD/SSD zbiorczo oraz dla każdej z perspektyw BSC są zawarte pod linkiem [www 3].
- Równania matematyczne opisujące zależności dynamiczne w modelu SD/SSD są zamieszczone pod linkiem [www 4] (przykład<sup>26</sup> dla wariantu Średniego)<sup>27</sup>.
- W celu obserwacji/ śledzenia oraz sterowania (modyfikacji/ regulacji) wybranych parametrów układu za pomocą „nastawnych kołowych potencjometrów” stworzono także Dashboard<sup>28</sup> (identyczny dla każdego wariantu) – link [www 5] (np. dla wariantu Średniego).



**Rys. 2.** Model SD/SSD – BSC. Perspektywy modelu oraz główne interakcje pomiędzy obszarami perspektyw BSC

Źródło: [www 1].

<sup>26</sup> Z użyciem danych modelu SSD WINGS odnoszącym się do wariantu Średniego.

<sup>27</sup> Uogólnione równania bez specyfiki WINGS EP21 do EP1 w zależności od wariantu (co odpowiada wartościom 0.0 w wektorach danych dla  $t$  od 0 do 60) przedstawiono pod linkiem [www 7].

<sup>28</sup> Oprócz Dashboardu zmiany (modyfikacje parametrów) mogą się odbywać poprzez „wpisanie” danych bezpośrednio do „wnętrza modelu” (równań matematycznych opisujących zależności pomiędzy elementami modelu).





### 1.2.2. Symulacje

Wykonano symulacje<sup>29</sup>:

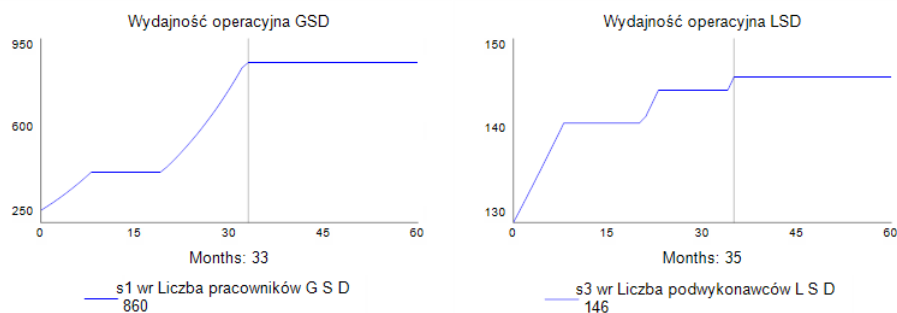
- W celu poprawy jakości dostarczania usług (w wydajności dostarczania występowała ewidentna „dziura” na szczycie – zaznaczona wielokątami na rys. 6<sup>30</sup>) z poziomu około 60% dla najbardziej niekorzystnego momentu (100% minus odczyt około 40% jako: 0.383 GSD oraz 0.404 LSD z wykresów – rys. 5) zwiększono w drodze symulacji początkową (dla chwili  $t = 0$ , czyli od początku dostarczania usług) liczbę zatrudnionych pracowników GSD z 200 do 296 oraz LSD ze 100 do 130, a także cenę za pakiet EP (150k€ do 215k€)<sup>31</sup> – co spowodowało wzrost jakości dostarczania usług do wymaganego poziomu około 80% (GSD oraz LSD: odczyt jako 0.19 oraz 0.188) w porównaniu z 60% poprzednio, a także zawężenie okna czasowego (wykresy) wskaźnika zaległości w dostarczaniu usług. W obydwu przypadkach (dla wariantu Średniego maksymalna liczba sprzedanych EP była oczywiście ta sama na poziomie 819 w miesiącu 28).
- Następnie zwiększono wydatki na sprzedaż i marketing z 5k€ na miesiąc do 10k€, co spowodowało wzrost maksymalnej liczby sprzedanych EP do poziomu 836 (również w miesiącu 28), a w wyniku tego wzrostu wzrosła również liczba dostarczanych EP przy tej samej (początkowej, jak wymieniono powyżej) liczbie pracowników GSD oraz podwykonawców LSD, co zaowocowało zmniejszeniem poziomu jakości dostarczanych usług o około 1% dla LSD (czyli zaległości wzrosły o 1%) i utrzymaniem podobnego poziomu jakości dla GSD.
- Dalej, zwiększenie efektywności zespołów sprzedażowych o 60% (co może zaistnieć w przypadku np. przejmowania biznesu 5G z uwagi na restrykcje nałożone na firmę Huawei) spowodowało wzrost liczby sprzedanych EP do 1150, co finalnie przyczyniło się do spadku jakości dostarczania usług do wyjściowego poziomu około 60% z początku symulacji.

<sup>29</sup> Ze względu na liczbę rysunków przypadających na każdy wariant ilustrujących przebiegi czasowe oraz reakcje na zmiany nie jest możliwe przedstawienie wszystkich realizacji w rozdziale, dlatego też dostrojenie modelu zostało zilustrowane na przykładzie wariantu Średniego dla wybranej grupy parametrów.

<sup>30</sup> Rysunek ten jest przeglądem symulacyjnym wybranych parametrów wydajności finansowej i operacyjnej.

<sup>31</sup> Jest zjawiskiem naturalnym, iż wzrost jakości to większy koszt dla firmy, a co za tym idzie większa cena za pakiet EP dla klienta.

- Powyższe zmiany, w odniesieniu do liczby zatrudnionych, wywołały reakcje autoregulacyjne modelu odzwierciedlone w przebiegach zatrudnienia (przebiegi GSD i LSD na rys. 5<sup>32</sup>) w perspektywie czasu, wskazujące na wymagane decyzje menedżerskie w stosunku do działań, które muszą być zaplanowane i podjęte w celu alokacji dodatkowych pracowników i podwykonawców (co pozwoli zbudować z wyprzedzaniem tzw. demand management dla danego wariantu/ scenariusza).

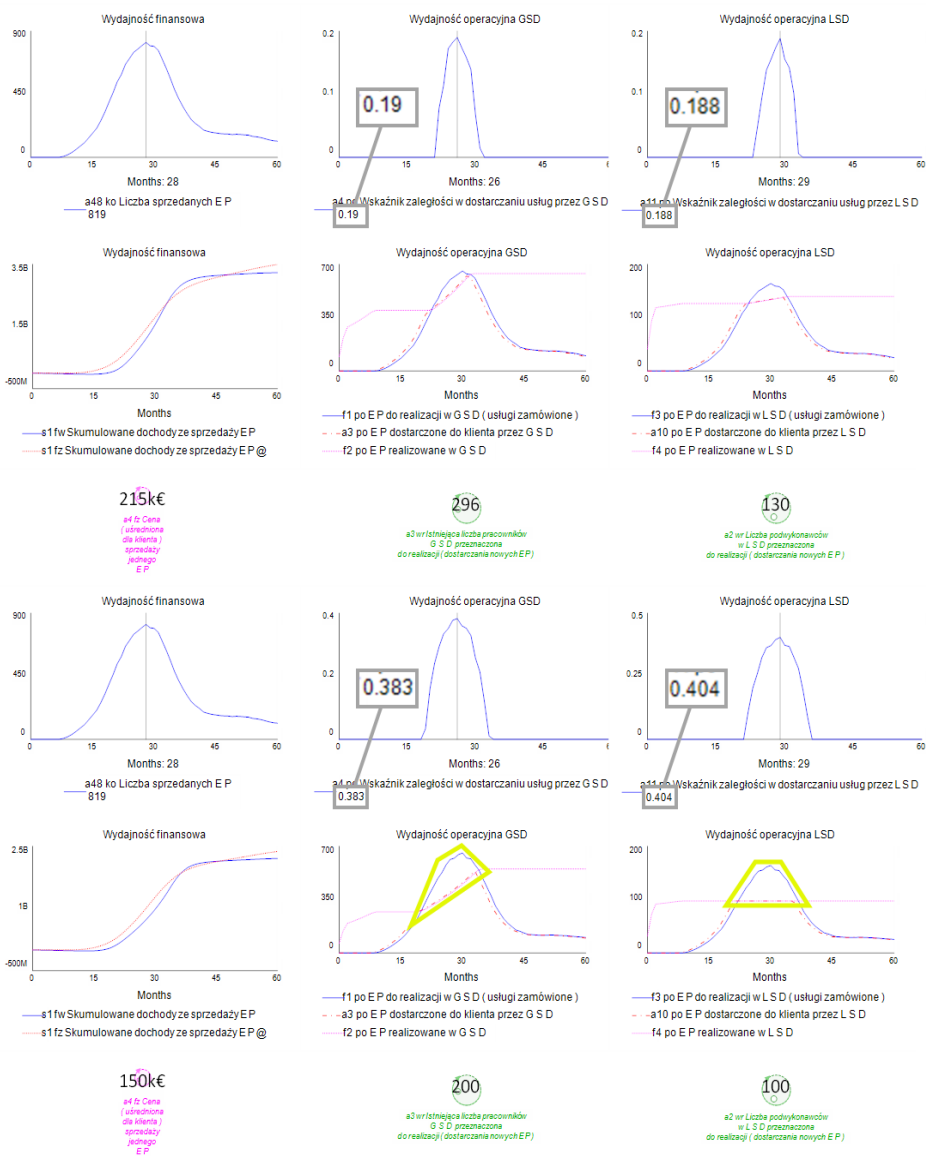


**Rys. 5.** Ewolucja zatrudnienia w czasie

- Podczas dalszych symulacji okazało się, że model SD/SSD wykazuje stosowną czułość na wszystkie adresowalne (możliwe do zmiany w warunkach rzeczywistych, gdyż takie były testowane) parametry<sup>33</sup>, dlatego też dostrajanie jest dość pracochłonnym procesem iteracyjnym.

<sup>32</sup> Rysunek ten przedstawia zmiany w liczbie niezbędnych pracowników i podwykonawców w perspektywie czasu 60 miesięcy, w wyniku zmiany początkowych (w chwili  $t = 0$ ) parametrów operacyjnych według przykładu dla wariantu Średniego.

<sup>33</sup> Nie zakłada się zmiany struktury modelu, co wiązałoby się ze zmianą konkretnego procesu, np. sprzedaży w kierunku platformowej, tj. bardziej zdigitalizowanej, ze skróceniem cyklu sprzedażowego z obecnych 6-9 miesięcy (cykl ten jest odzwierciedlony jako czynnik losowy, również o rozkładzie równomiernym prawdopodobieństwa, opóźnienia w równaniu 7) do znacznie krótszego.



**Rys. 6.** Wariant Średni – wykresy symulacyjne

Aby usystematyzować dostrajanie (lub regulację) opracowano funkcję  $\Delta(s1\_fz - s1\_fw)$ :

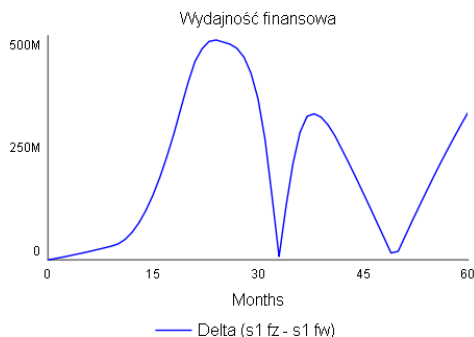
$$\Delta(s1\_fz - s1\_fw) = MIN\{AVERAGE\{ABSs1\_fz(t) - s1\_fw(t)\}\} \quad (1)$$

gdzie  $s1\_fz(t)$  (rys. 9<sup>34</sup>) to skumulowane dochody ze sprzedaży EP @ (przy „stałej cenie”), a  $s1\_fw(t)$  (rys. 8<sup>35</sup>) – skumulowane dochody ze sprzedaży EP (przy „cenie zmiennej”).

$$s1\_fz(t) = s1\_fz * (t - \Delta t) + (f1\_fz - f2\_fz) * \Delta t$$

$$s1\_fw(t) = s1\_fw * (t - \Delta t) + (f1\_fw - f2\_fw) * \Delta t$$

Dla funkcji  $\Delta(s1\_fz - s1\_fw)$  ustalono eksperymentalnie na podstawie zachowania się modelu, iż powinna zostać przeprowadzona minimalizacja średniej  $\Delta(s1\_fz - s1\_fw)$  w ujęciu bezwzględnym, czyli „wypłaszczenie” amplitudy (w mln EUR) przebiegu funkcji  $\Delta(s1\_fz - s1\_fw)$ , co w symulacjach będzie odzwierciedlone jako dążenie do niezachowania (jeśli  $s1\_fz > s1\_fw$ ) symetrii (i amplitudy) wierzchołków funkcji  $\Delta(s1\_fz - s1\_fw)$ , aby model SD/SSD osiągał wymagane parametry operacyjne według założonych kryteriów biznesowych. Przykładowy przebieg funkcji  $\Delta(s1\_fz - s1\_fw)$  przedstawiono na rys. 7<sup>36</sup>.



**Rys. 7.** Eksperymentalnie dobrana funkcja jakości symulacji  $\Delta(s1\_fz - s1\_fw)$

<sup>34</sup> Rysunek ten przedstawia Perspektywę Finansową widzianą według parametrów sprzedaży „do klienta” (tzw. ujęcie zewnętrzne).

<sup>35</sup> Rysunek ten przedstawia Perspektywę Finansową widzianą według wewnętrznych parametrów finansowych (dla użytku kalkulacji wewnętrznych – tzw. ujęcie wewnętrzne).

<sup>36</sup> Jest to wykres syntetycznej funkcji  $\Delta$  opracowanej na potrzeby monitorowania jakości działania modelu SD/SSD według postawionych celów biznesowych (patrz wprowadzenie do rozdziału).

Zagadnienie „ceny zmiennej” (ilustracja na rys. 9 – parametr  $a7\_fw$ ) pojawia się z uwagi na zmienną liczbę sprzedanych, a dostarczonych usług EP w czasie  $t$ . Tematyka ta jest realizowana na potrzeby tzw. rachunkowości wewnętrznej w firmie (model cenowy użyty w SD/SSD, jak również w rzeczywistości operacyjnej firmy to „cost+”, stąd w zależności od liczby sprzedanych EP występuje rzeczona „cena zmienna”). Ilustrację poprzez parametry otoczone pomarańczowymi „owalami” w celu wyjaśnienia „ceny zmiennej” przedstawiono na rys. 7, a równania (przykładowe dla GSD) odzwierciedlające jej metodykę zamieszczono poniżej.

Aby wyznaczyć „cenę zmienną” w danej w chwili  $t$ , oblicza się sumaryczne:

- „Przychody ze wszystkich EP miesiąc do miesiąca” (przepływ  $f1\_fw$  (2)) na podstawie:
  - „Wartości finansowej wszystkich EP (cost+) GSD”  $a16\_fw$  (4) oraz
  - „Wartości finansowej wszystkich EP (cost+) LSD”  $a12\_fw$ .
  - \* „Wartość finansowa wszystkich EP (cost+) GSD” zawiera w sobie koszty EP ( $a17\_fw$  (5)), do których dodawana jest zawsze ustalona marża ( $a15\_fw$  (3)).
  - \* Analogicznie zachodzi dla „Wartość finansowa wszystkich EP (cost+) LSD – według zasad tzw. modelu cenowego „cost+”.
- Następnie „Przychody ze wszystkich EP miesiąc do miesiąca” ( $f1\_fw$ ) są dzielone przez „Liczbę sprzedanych EP” ( $a48\_ko$  (6)) w chwili  $t$ , co daje przychód przypadający na EP, i ta wielkość byłaby hipotetyczną „ceną zmienną” ( $a7\_fw$  (8)) do zaoferowania klientowi w danej chwili  $t$ , gdyż wartość jej zmienia się (w ujęciu miesięcznym) z uwagi na liczbę sprzedanych i dostarczonych usług EP w chwili  $t$ .
- Realizacja ww. odbywa się w Perspektywie FINANSOWEJ (rys. 8) – ujęcie od wewnątrz (tzw. rachunkowość wewnętrzna). Przykład równań dla GSD:

$$f1\_fw = a16\_fw + a12\_fw \quad (2)$$

$$a15\_fw = 0.5 \quad (3)$$

$$a16\_fw = (a17\_fw * a3\_po) * (1 + a15\_fw) \quad (4)$$

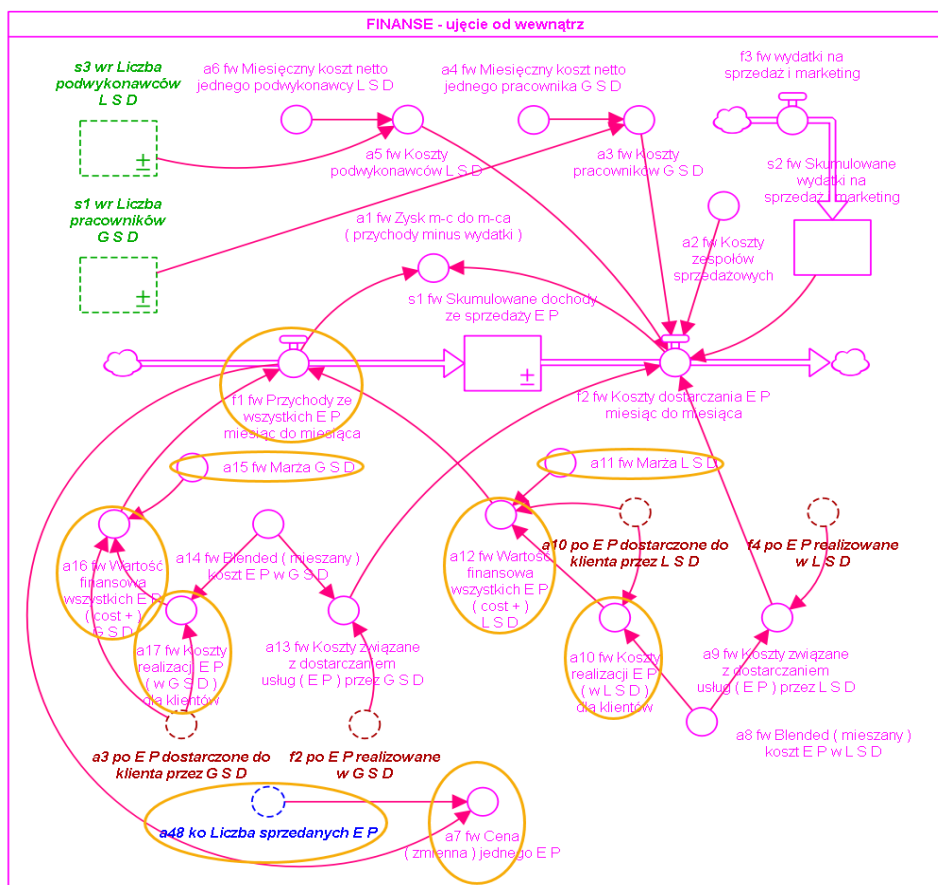
$$a17\_fw = a3\_po * a14\_fw \quad (5)$$

$$a48\_ko = \sum_{m=1}^{21} w_m\_ko * a2m\_ko \quad (6)$$

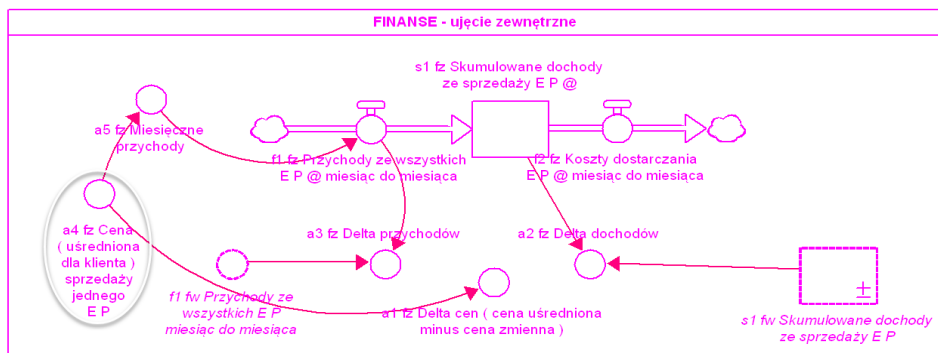
będący sumą wartości 21 funkcji w danej chwili  $t$ , które to funkcje pochodzą z 21 modułów WINGS SSD (rys. 2). Przykładowa funkcja  $w1\_ko.a2\_ko$ :

$$w1\_ko.a2\_ko = DELAY(a1\_ko, RANDOM(6,9),0) \quad (7)$$

$$a7\_fw = f1\_fw / a48\_ko \quad (8)$$

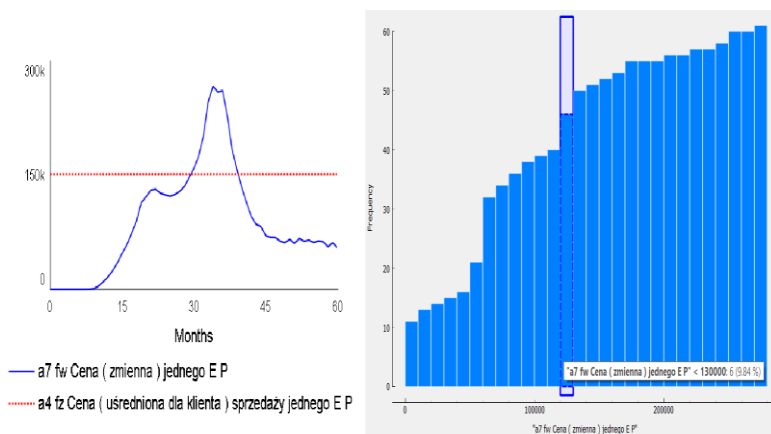


**Rys. 8.** Perspektywa FINANSOWA (część rys. 3) – ujęcie od wewnątrz  
(na rysunku nie są ujęte połączenia do innej perspektywy BSC)



**Rys. 9.** Perspektywa FINANSOWA (część rys. 3) – ujęcie od zewnątrz  
(na rysunku nie są ujęte połączenia do innej perspektywy BSC)

„Cena stała”<sup>37</sup> (parametr a4\_fz Cena (uśredniona dla klienta) sprzedaży jednego EP – rys. 8 (zaznaczenie) oraz ujęcie w porównaniu do „ceny zmiennej” – rys. 9<sup>38</sup>) to cena obowiązująca w całym okresie 5 lat dostarczania usług (klienti oczekują przewidywalności ceny z uwagi na TCO – Total Cost of Ownership [Lamalle, 2014]). W idealnym przypadku skumulowane wartości „ceny zmiennej” i „ceny stałej” (przebiegi na wykresach) w okresie 5 lat powinny się pokrywać (skumulowane wartości s1\_fz i s1\_fw – rys. 6 – patrz rys. pod nazwą Wydajność finansowa), gdyż wyższa skumulowana wartość „ceny stałej” powodowałaby spadek konkurencyjności<sup>39</sup>, niższa natomiast od skumulowanej wartości „ceny zmiennej” powodowałaby nieosiągnięcie założonej rentowności dostarczania usług, a to mogłoby spowodować wymknienie się spod kontroli i groziłoby negatywnymi wynikami finansowymi (koszty większe od zysków).



**Rys. 10.** Perspektywa FINANSOWA: przebiegi ceny „zmiennej” i ceny „stałej”

<sup>37</sup> Cena stała jest ceną referencyjną w procesie sprzedaży, gdzie w zależności od rynku i klienta (i negocjacji cenowych) jej wartość może ulegać wahaniom (powinna być większa, lecz nie może być generalnie niższa niż cena referencyjna, jednak nie zawsze jest to możliwe, również z uwagi na sprzedaż związaną sprzętu i usług). Wartość ceny referencyjnej powinna być dostosowaną wartością w celu minimalizacji delty między dochodami fl\_fz i fl\_fw (użyto dochodów skumulowanych) i różni się ona od miary przeciętnej/pozycyjnej (największej częstości występowania poziomu ceny zmiennej) – rys. 9 (prawy wykres), która nie zapewniła minimalizacji wspomnianej delty (lecz została użyta jako dość dobre przybliżenie na początku strojenia modelu (dla wariantu Średniego).

<sup>38</sup> Część „lewa” rysunku: „Cena zmienna” i „Cena stała” (dla przykładu); część „prawa” rysunku: to największa częstość występowania danego poziomu ceny (zmiennej) na przestrzeni 60 miesięcy i jej wartość.

<sup>39</sup> Należy zauważyć, że firma działa na rynku oligopolistycznym, gdzie wszyscy dostawcy (Ericsson, Huawei, Nokia, Samsung, ZTE itd.) mają obecnie porównywalne koszty dostarczania usług.

W wyniku przeprowadzonych symulacji ustalono również sposób dostrajania modelu SD/SSD (punkty 1-3, ewentualnie 4) dla wybranych/głównych parametrów operacyjnych:

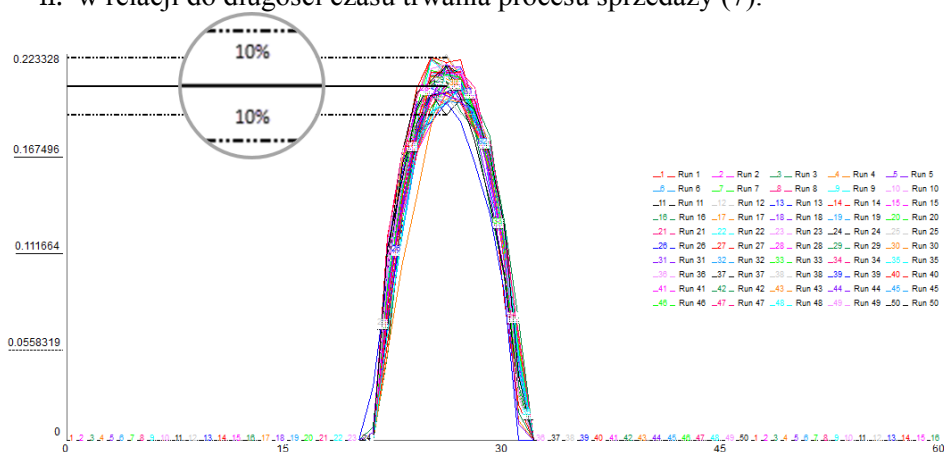
1. Zmiany liczby zatrudnionych w GSD i LSD powinny być tak dobrane, aby wskaźniki jakości dostarczania usług były „w szczycie” na poziomie około 80% każdy (nie powinna istnieć asymetria w wartościach). Wskaźnik zaległości w dostarczaniu usług wynosi wtedy około 20% dla GSD i LSD w najbardziej niekorzystnym momencie (rys. 6).
2. Następnie należy dobrać tzw. cenę stałą za EP według wykresu na rys. 7, z zachowaniem dostrajanej ceny według przebiegu (kształtu) funkcji  $\Delta(s1_{fz} - s1_{fw})$  i ewentualnym wypłaszczeniem amplitudy przebiegu tej funkcji (to ostanie może zostać dokonane np. poprzez zmniejszenie marży GSD np. z 50% do wartości mniejszych, porównywalnych z marżą LSD, w celu poprawy konkurencyjności ceny EP), lecz z zachowaniem jej asymetrii i proporcji wierzchołków.
 

(a) Przeszacowywanie wartości ceny stałej<sup>40</sup> prowadzi do zmniejszenia konkurencyjności usług EP (bez poprawy wskaźników jakości dostarczania usług), zatem zalecany jest sposób postępowania opisany powyżej.
3. Kroki 1 i 2 odzwierciedlają dostrojenie wystarczające do ogólnej wstępnej oceny symulacyjnej działania modelu SD/SSD i prezentacji otrzymanych rezultatów (benchmarkingu), w którym należy wziąć pod uwagę możliwy margines błędu (odchylenie) otrzymanych wyników od wartości średnich nastaw w danym wariancie (punkt 4 poniżej). Jeśli niezbędna jest ocena bardziej szczegółowa i bardziej precyzyjne dostrojenie, należy wykonać krok 4, niemniej wstępnie można założyć, że margines błędu to  $\approx 5\%$  od wyliczonych wartości<sup>41</sup>.

<sup>40</sup> W celu ilustracji opisowej/praktycznej w kontekście ceny stałej przyjęto, że cena dla klienta za jeden EP to 150 000 € i jest niezmienna w każdym miesiącu w całym okresie (60 miesięcy) dostarczania usług dla wszystkich klientów globalnie. W idealistycznej sytuacji (wariant Idealny), odchodząc teraz nieco od wariantu Średniego, może to być 21 EP \* 150 000 € = 3 150 000 €, które zapłaci jeden klient jako opłatę jednorazową w chwili  $t$  w momencie dostarczenia usług (z możliwością używania pakietu EP przez okres 5 lat), jeśliby zakupił wszystkie EP w pełnym wymiarze (czyli cały oferowany pakiet składający się z 21 EP). Powyższe obowiązuje także wtedy, gdy poszczególne EP nie są zakupione w 100-procentowym wymiarze przy wyborze pakietu EP. Klienci jednak wiedzą, że 21 Elementów Portfolio tworzy komplementarny zestaw usług, stąd przewiduje się raczej niewielki odsetek rezygnacji z poszczególnych EP, co jest odzwierciedlone w wynikach wariantów od Antyidealnego do Idealnego.

<sup>41</sup> Jak już wspomniano w podrozdziale 1.1.1, maksymalna średnia wartość błędu dla  $dt = 1$  wynosiła  $\approx 5\%$  w ogóle dla przetestowanych różnych grup parametrów (średnio dla magazynów, przepływów i konwerterów).

4. W celu uzyskania możliwie jak najbardziej dokładnych nastaw modelu należałoby powtórzyć symulacje 1 i 2 około 50 razy (eksperymentalnie taka liczność próby została dobrana jako reprezentatywna w próbie ogólnej wynoszącej 100 powtórzeń). Na przykład w celu dokładnej kontroli osiąganego jakości dostarczania usług w chwili  $t$  w 50 powtórzeniach pojawia się wartość MIN oraz MAX wskaźnika jakości dostarczania usług – i na tej podstawie można by wyznaczyć wartość średnią korytarza. W wyniku prób wykonanych przy doborze  $dt$  (wspomniane w podrozdziale 1.1.1, włączając w to dokładność całkowania) oraz analizując wyniki symulacyjne – korytarz odchylenia od wartości średniej to w przybliżeniu  $\pm 10\%$  dla wskaźnika zaległości (GSD i LSD) dostarczania usług (w połowie – miesiąc 29) trwania okresu symulacji<sup>42</sup> (przykładowy rys. 11<sup>43</sup>).
- (a) Potrzeba powtórzeń symulacji w celu oszacowania odchylenia bierze się głównie z faktu występowania dwóch czynników losowych, które zmieniają charakter przebiegów – powodują rozrzut wartości (występują one w procesie sprzedaży – Perspektywa Klienta BSC):
- w części WINGS odnoszącej się do wskaźnika nominalnego WINGS EP (dla każdego wariantu); (charakter losowości wyjaśniony w podrozdziale 1.1.1),
  - w relacji do długości czasu trwania procesu sprzedaży (7).



Rys. 11. Przebieg wskaźnika zaległości dostarczania usług dla GSD

<sup>42</sup> Odchylenie w symulacjach maleje „z czasem”. Dla zastosowanego w modelu SD/SSD  $dt = 1$  wartość procentowa w miesiącu 60. nie przekraczała 1% odchylenia od mediany wartości średnich wszystkich przetestowanych  $dt$  ( $dt = 1, dt = 0.25, dt = 0.1, dt = 0.03$ ) dla większości grup parametrów (magazyn, przepływ, parametr pomocniczy).

<sup>43</sup> Jest to odwzorowanie zachowania parametru  $a4\_po$  w perspektywie Procesy Wewnętrzne na rys. 3 w 50 powtórzeniach (Run 1 do Run 50).

### 1.2.3. Warianty

Do zaprojektowania możliwych scenariuszy rozwojowych w okresie 5-letnim (dla procesu sprzedaży i dostarczania EP) w celu wsparcia decyzji menedżerskich wykonano modelowanie (stworzono 5 głównych wersji modelu SD/SSD dla wariantów Antyidealny, Słaby, Średni, Dobry, Idealny z oddzielnymi zestawami danych [www 8]) oraz przeprowadzono symulacje dla każdego wariantu według sposobu dostrajania modelu SD/SSD, punkty 1-3 – opisane w poprzedniej części rozdziału. Zbiorcze rezultaty zamieszczono na rys. 12<sup>44</sup>.

Do wygenerowanej w drodze symulacji przez model SD/SSD liczby sprzedanych i dostarczonych EP (parametr  $a48\_ko$ ) „w szczycie”, czyli miesiącu, w którym wystąpi ekstremalne „obciążenie” organizacji firmy ( $\approx$  w połowie okresu 5-letniego), dostrajano parametry operacyjne  $a3\_wr$  oraz  $a2\_wr$ , aby osiągnąć wymaganą całkowitą jakość i nie mniejszą niż 80% w najniekorzystniejszej chwili czasu  $t$  (dla uproszczenia – na potrzeby opisu w rozdziale zachowano wszystkie pozostałe parametry modelu bez zmian) – wyniki przedstawia tabela 2<sup>45</sup>. W następnych krokach badano wpływ marży na osiągany dochód przy różnych poziomach „ceny stałej” dla klientów (dostrajanie według funkcji  $\Delta(s1\_fz - s1\_fw)$  (1)), gdzie przetestowano zakres marż dla GSD (rys. 11) od minimalnej dopuszczalnej (10%) do maksymalnej (50%) z  $\approx$  5-procentowym średnim marginesem błędu wyniku (według punktu 3 dostrajania modelu SD/SSD), podczas gdy marża LSD była cały czas na niezmiennym poziomie 30%.

Okazało się również, że warianty Dobry i Idealny „zachodzą” na siebie w korytarzu 5-procentowego marginesu błędu wyników (pierwotną przyczyną jest fakt, że dla wariantów Dobrego i Idealnego, kiedy wartość wylosowana przekraczała 1, wartość ta była zastępowana liczbą 1 jako maksymalnym możliwym poziomem sukcesu – podrozdział 1.1.1).

Podsumowujące wyniki przedstawiające niezbędne parametry operacyjne: zatrudnienie początkowe oraz jego ewolucję w czasie<sup>46</sup> przedstawiono na rys. 13<sup>47</sup>, a co za tym idzie wsparcie decyzyjne w zależności od realizowanego scenariusza według bieżącej/rzeczywistej sytuacji.

<sup>44</sup> Na rysunku tym przedstawiono wyniki symulacji w dostrojeniu wariantów: Antyidealny, Słaby, Średni, Dobry, Idealny (w różnych opcjach marży i ceny uśrednionej za EP dla klienta). Rysunek ten odzwierciedla także podobieństwa w zachowaniu każdego wariantu, gdzie dominanta zachowania jest uzależniona od procesu SSD.

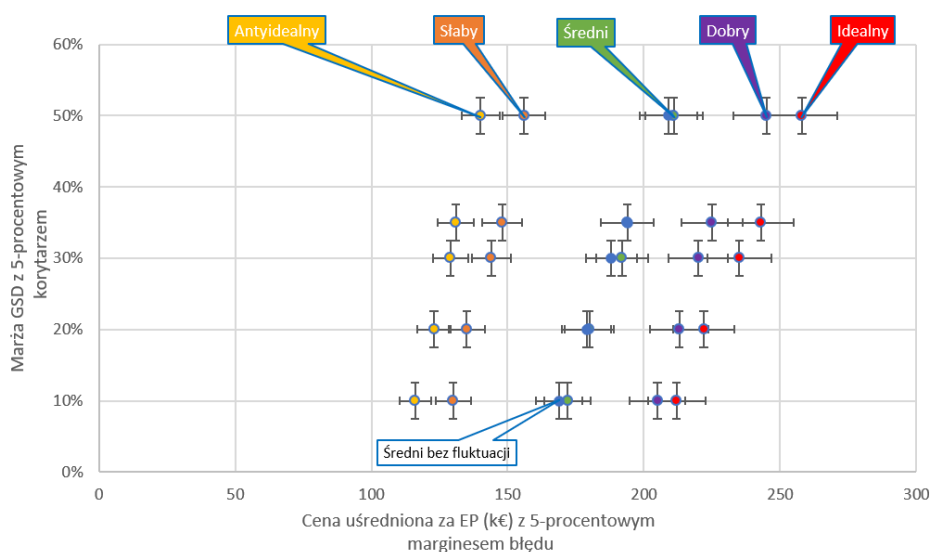
<sup>45</sup> Wyniki dotyczące zatrudnienia są odzwierciedlone dla chwili  $t = 0$  (usługi jeszcze nie są dostarczane). Liczba dostarczonych EP wyraża uzyskane wielkości według wymagań kompromisu jakości w najbardziej niekorzystnym momencie dostarczania usług.

<sup>46</sup> Tak zwane podstawowe benchmarki operacyjne.

<sup>47</sup> Przy początkowym zatrudnieniu według tabeli 2.

**Tabela 2.** Dostrojone podstawowe parametry operacyjne – sprzedane EP oraz zatrudnienie – w modelu SD/SSD

| Wariant     | a48_ko: liczba sprzedanych i dostarczonych EP „w szczycie” | a3_wr: początkowa liczba pracowników GSD ( $t = 0$ ) | a2_wr: początkowa liczba podwykonawców LSD ( $t = 0$ ) |
|-------------|------------------------------------------------------------|------------------------------------------------------|--------------------------------------------------------|
| Idealny     | 1000                                                       | 353                                                  | 155                                                    |
| Dobry       | 960                                                        | 343                                                  | 145                                                    |
| Średni      | 810                                                        | 286                                                  | 127                                                    |
| Słaby       | 630                                                        | 213                                                  | 103                                                    |
| Antyidealny | 570                                                        | 199                                                  | 89                                                     |

**Rys. 12.** Graficzna reprezentacja pozycjonowania dostrojonych wszystkich wariantów w ujęciu marży GSD do ceny za EP

Źródło: [www 9] (zbiorcza tabela danych).



### 1.2.4. Wnioski

W hybrydowym modelu SD/SSD połączono razem operacyjne dane ilościowe gromadzone przez firmę w strukturę SD oraz dane jakościowe w strukturę SSD z użyciem metody WINGS – z procesu, który nie mógł być opisany w ramowej strukturze ilościowej SD (wiedza o samych elementach i ich interakcjach była na tyle ogólna, iż ich jawne relacje funkcjonalne nie mogły być określone lub sformalizowane). Podejście to w środowisku BSC stworzyło zupełnie nowy pod względem metodycznym sposób łączenia danych ilościowych i jakościowych z wykorzystaniem technik obliczeniowych (dyskretyzacji danych procesów ciągłych). Jednak należy pamiętać, iż **każdy przypadek jest indywidualny w zależności od rozpatrywanego problemu (nie ma jednolitego wzorca hybrydyzacji)**, stąd struktura metodyczna przedstawiona w rozdziale może być zreplikowana, lecz projekt modelu/układu SD/SSD musi być stworzony dla konkretnego przypadku biznesowego (**co, z czym łączyć**). Dla postawionego zadania wprowadzenia zaawansowanych usług sieci 5G i przygotowania organizacji pod 6G zostało zrealizowane:

- dynamiczne oszacowanie liczby pozyskanych klientów w obszarze akwizycji,
- konwersja na późniejszą sprzedaż z uwzględnieniem jej specyfiki,
- dynamiczne zwymiarowanie organizacji dostarczania usług 5G<sup>48</sup>:
  - w relacji do spodziewanej absorpcji tychże usług przez operatorów sieci telekomunikacyjnych,
  - z pierwszorzędnym celem osiągnięcia dochodowości firmy na założonym poziomie (minimalna stała marża GSD 10%, maksymalna 50%),
  - z zapewnieniem akceptowalnego kompromisu jakości<sup>49</sup> przekładającego się w praktyce na poziom zaległości w realizacji założonych funkcjonalności usług EP, w najbardziej niekorzystnej chwili czasu, a wpływającego na koszty własne.

Koszty opracowania modelu to zaledwie niewielka część z wygenerowanych korzyści (znane firmy konsultingowe Bain, Gartner artykułują korzyści wynikające ze wsparcia strategicznych decyzji menedżerskich na poziomie 4-6% osiąganey marży<sup>50</sup>). Podsumowując, w projektowaniu możliwych scenariuszy nie chodziło o to, aby przewidzieć, który scenariusz jest właściwy, ale

<sup>48</sup> Wraz z przygotowaniem kadry do 6G od drugiej połowy okresu dostarczania usług 5G.

<sup>49</sup> Według zależności empirycznej oszacowanej w firmie jako wskaźnik ewentualnego zmniejszenia liczby klientów z uwagi na działalność operacyjną.

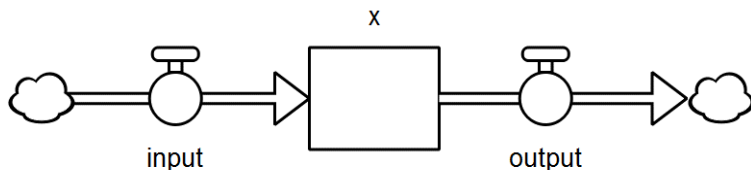
<sup>50</sup> Jest to ich praktyka obserwacyjna (wiedza ekspercka), niepoparta jednakże pozycją naukową.

o zdefiniowanie wachlarza możliwości, cech, które je wyróżniają, a najlepiej, jaki kierunek działań (oraz konkretnie jakie parametry operacyjne w zależności od ewolucji rzeczywistej sytuacji) umożliwi przedsiębiorstwu odniesienie sukcesu niezależnie od tego, który ostatecznie scenariusz zostanie zrealizowany.

## 1.3. Dodatek

### 1.3.1. SD

Do opisu podstawowych zależności matematycznych (równania: 11-12) wykorzystano równania różnicowe (w implementacji numerycznej tej metody) z przyjętą notacją przyrostową  $\Delta t$  (równania: 9-10), gdzie dla uogólnionego modelu magazynu  $x$  (na poniższym rys. 14) występują również zawory, które oznaczają sterowanie przepływami, a także konwertery<sup>51</sup>. Pojemniki w otoczeniu systemu, oznaczone jako „chmurki”, reprezentują nieograniczone źródła zasobów dla wpływów do zbiornika w systemie lub nieograniczone pojemniki dla wpływów ze zbiornika w systemie.



**Rys. 14.** Magazyn ( $x$ ), wpływ (input) i wypływ (output)

$$x(t) = x(t - \Delta t) + (\text{input} - \text{output})\Delta t \quad (9)$$

gdzie  $\Delta t$  – przedział czasu.

$$\frac{x(t) - x(t - \Delta t)}{\Delta t} = (\text{input} - \text{output}) \quad (10)$$

Równania te to inna postać przybliżonych równań dyskretnych, które wynikają z równań dla ciągłego czasu (postać całkowa i różniczkowa):

<sup>51</sup> Nieoznaczone na schematycznym rys. 14. Funkcje: przechowywanie wartości stałych, definiowanie zewnętrznych danych wejściowych do modelu (użyte do importu danych z SSD WINGS dla „Wskaźnika sukcesu ofertowego – sprzedaży EP (WINGS EP)” – patrz również podrozdział 1.3.2, gdzie również jest mowa o wskaźniku WINGS EP), obliczanie relacji algebraicznych, repozytorium funkcji graficznych.

$$x(t) = \int (\text{input} - \text{output})dt \quad (11)$$

$$\frac{dx(t)}{dt} = \text{input} - \text{output} \quad (12)$$

### 1.3.2. SSD

Do opisu podstawowych zależności matematycznych służą równania (13-16). W metodzie tej ocenie podlegają dwie charakterystyki dotyczące roli danego składnika w systemie: siła (znaczenie) danego składnika w systemie oraz siła, z jaką oddziałuje on na inne składniki systemu. Wszystkie wartości określone przez użytkownika wprowadza się do macierzy **D** bezpośredniego znaczenia/wpływów. Macierz ta jest skalowana według formuły:

$$\mathbf{S} = \frac{1}{s} \mathbf{D} \quad (13)$$

gdzie **S** jest *skalowaną macierzą siły wpływów*, a czynnik skalujący jest zdefiniowany jako suma elementów macierzy **D**:

$$s = \sum_{i=1}^n \sum_{j=1}^n |d_{ij}| \quad (14)$$

Macierz całkowitego znaczenia/wpływów **T** oblicza się według wzoru:

$$\mathbf{T}(t) = \mathbf{S} + \mathbf{S}^2 + \mathbf{S}^3 + \dots + \mathbf{S}^t \quad (15)$$

gdzie  $t = 0, 1, \dots, 60$ . Następnie oblicza się całkowitą podatność  $R_i$  w celu uzyskania wielkości obrazującej łączne wpływy wywierane przez składniki systemu na wybrany składnik  $i$ . Jeżeli zsumuje się elementy należące do kolumny  $i$  w macierzy **T**, uzyska się taką wielkość  $R_i$ . W kontekście części jakościowej modelu hybrydowego będzie to wskaźnik nominalny WINGS EP, a w ujęciu dynamicznym [Adamus-Matuszyńska, Michnik, Polok, 2019] wektor 61 wielkości nominalnych WINGS EP<sup>52</sup> od  $t = 0$  do  $t = 60$  (okres modelowania 5 lat w ujęciu miesięcznym – występuje w podrozdziale 1.2):

$$R_i = \sum_{j=1}^n t_{ji}(t) \quad (16)$$

gdzie  $t_{ji}(t) = [\mathbf{T}(t)]_{ji}$ .

<sup>52</sup> Odwołanie do wielkości nominalnych WINGS EP pojawia się także w podrozdziale 1.1.1.

## Literatura

- Adamus-Matuszyńska A., Michnik J., Polok G. (2019), *A Systemic Approach to City Image Building. The Case of Katowice City*, "Sustainability", Vol. 11(16), 4470, Publisher: Multidisciplinary Digital Publishing Institute, s. 1-20, <https://www.mdpi.com/2071-1050/11/16/4470> (dostęp: 16.03.2021).
- Akkermans H.A., Oorschot K.E. van (2002), *Developing a Balanced Scorecard with System Dynamics*, Proceedings of the 20th International Conference of the System Dynamics Society, July 28 – August 1, at Palermo, Italy, <https://research.tue.nl/en/publications/developing-a-balanced-scorecard-with-system-dynamics> (dostęp: 17.03.2021).
- Axelos (2011), *Glosariusz ITIL® wraz ze skrótami*, [https://www.axelos.com/corporate/media/files/glossaries/itil\\_2011\\_glossary\\_pl-v1-0.pdf](https://www.axelos.com/corporate/media/files/glossaries/itil_2011_glossary_pl-v1-0.pdf) (dostęp: 25.07.2021).
- Banaś D., Michnik J. (2019), *Evaluation of the Impact of Strategic Offers on the Financial and Strategic Health of the Company – A Soft System Dynamics Approach*, "Mathematics", Vol. 7(2), s. 208.
- Barnabè F. (2011), *A System Dynamics-based Balanced Scorecard to Support Strategic Decision Making: Insights from a Case Study*, "International Journal of Productivity and Performance Management", Vol. 60(5) s. 446-473.
- Checkland P. (2000), *Soft Systems Methodology: A Thirty Year Retrospective*, "Systems Research and Behavioral Science", Vol. 17(S1), S11-S58.
- Coyle G. (2000), *Qualitative and Quantitative Modelling in System Dynamics: Some Research Questions*, "System Dynamics Review", Vol. 16(3), s. 225-244.
- Forrester J.W. (2013), *Industrial Dynamics*, Martino Fine Books, Eastford, CT, USA.
- Gensun F. (1996), *Whittaker-Kotelnikov-Shannon Sampling Theorem and Aliasing Error*, "Journal of Approximation Theory", Vol. 85(2), s. 115-131.
- Hanafizadeh P., Aliehyaei R. (2011), *The Application of Fuzzy Cognitive Map in Soft System Methodology*, "Systemic Practice and Action Research", Vol. 24(4), s. 325-354.
- Kaplan R.S., David P.N. (1992), *The Balanced Scorecard – Measures that Drive Performance*, "Harvard Business Review", Vol. 70(1), s. 71-79.
- Kasperska E. (2005), *Modelling Embedded in Learning – The Acceleration of Learning by the Use of the Hybrid Models on the Base of System Dynamics Paradigm [w:] Systemy wspomagania organizacji SWO' 2005*, Akademia Ekonomiczna, Katowice, s. 410-417.
- Kosko B. (1986), *Fuzzy Cognitive Maps*, "International Journal of Man-Machine Studies", Vol. 24(1), s. 65-75.
- Lamalle F. (2014), *Cloud Collaboration: What's the Impact on TCO an ROI?* Orange Business Services, <https://www.orange-business.com/en/blogs/enterprisingbusiness/collaboration/collaborating-in-the-cloud-what-s-the-impact-on-tco-and-roi> (dostęp: 7.04.2021).

- Lane D.C., Rogelio O. (1998), *The Greater Whole: Towards a Synthesis of System Dynamics and Soft Systems Methodology*, "European Journal of Operational Research", Vol. 107(1), s. 214-235.
- Linard K., Joseph Y. (2000), *The Dynamics of Organisational Performance Development of a Dynamic Balanced Scorecard*, [https://www.researchgate.net/publication/221503130\\_The\\_Dynamics\\_of\\_Organisational\\_Performance\\_Development\\_of\\_a\\_Dynamic\\_Balanced\\_Scorecard](https://www.researchgate.net/publication/221503130_The_Dynamics_of_Organisational_Performance_Development_of_a_Dynamic_Balanced_Scorecard) (dostęp: 17.03.2021).
- Luke H.D. (1999), *The Origins of the Sampling Theorem*, "IEEE Communications Magazine", Vol. 37(4), s. 106-108.
- Luna-Reyes L.F., Andersen D.L. (2003), *Collecting and Analyzing Qualitative Data for System Dynamics: Methods and Models*, "System Dynamics Review", Vol. 19(4), s. 271-296.
- Mendoza G.A., Prabhu R. (2006), *Participatory Modeling and Analysis for Sustainable Forest Management: Overview of Soft System Dynamics Models and Applications*, "Forest Policy and Economics", Vol. 9(2), s. 179-196.
- Michnik J. (2013), *Weighted Influence Non-linear Gauge System (WINGS) – An Analysis Method for the Systems of Interrelated Components*, "European Journal of Operational Research", Vol. 228(3), s. 536-544.
- Nielsen S., Nielsen E.H. (2015), *The Balanced Scorecard and the Strategic Learning Process: A System Dynamics Modeling Approach* [w:] *Advances in Decision Sciences 2015*, Publisher: Hindawi, e213758, <https://www.hindawi.com/journals/ads/2015/213758/> (dostęp: 17.03.2021).
- Rabelo L., Helal M., Jones A.T., Min H. (2005), *Enterprise Simulation: A Hybrid System Approach*, "International Journal of Computer Integrated Manufacturing", Vol. 18(6), s. 498-508.
- Rodriguez-Ulloa R., Paucar-Caceres A. (2005), *Soft System Dynamics Methodology (SSDM): Combining Soft Systems Methodology (SSM) and System Dynamics (SD)*, "Systemic Practice and Action Research", Vol. 18(3), s. 303-334.
- Rodríguez-Ulloa R.A., Montbrun A., Silvio Martínez-Vicente (2011), *Soft System Dynamics Methodology in Action: A Study of the Problem of Citizen Insecurity in an Argentinean Province*, "Systemic Practice and Action Research", Vol. 24(4), s. 275-323.
- Smith S.W. (2007), *Cyfrowe przetwarzanie sygnałów. Praktyczny poradnik dla inżynierów i naukowców*, Wydawnictwo BTC, Warszawa.
- Sterman J. (2000), *Business Dynamics: Systems Thinking and Modeling for a Complex World*, Irwin/McGraw-Hill, Boston.
- Tashakkori A., Teddlie C. (2010), *SAGE Handbook of Mixed Methods in Social & Behavioral Research*, 2455 Teller Road, Thousand Oaks California 91320 United States: SAGE Publications, Inc., <http://methods.sagepub.com/book/sage-handbook-of-mixed-methods-social-behavioral-research-2e> (dostęp: 20.03.2021).

- Vetter P., Harish V. (2021), *Nokia Brings its 6G Expertise to New US NextG Initiative*, <https://www.nokia.com/blog/nokia-brings-its-6g-expertise-to-new-us-nextg-initiative/> (dostęp: 5.05.2021).
- Zolfagharian M., Romme A.G.L., Walrave B. (2018), *Why, When, and How to Combine System Dynamics with Other Methods: Towards an Evidencebased Framework*, "Journal of Simulation", Vol. 12(2), s. 98-114.
- [www 1] [https://drive.google.com/file/d/1BfXD\\_x5NhIX3VlOOn2nw\\_wNVBqHplYm7/view?usp=sharingModel SD/SSD - BSC](https://drive.google.com/file/d/1BfXD_x5NhIX3VlOOn2nw_wNVBqHplYm7/view?usp=sharingModel%20SD/SSD-BSC).
- [www 2] [https://drive.google.com/drive/folders/1kmCOspV-hgP4K3lZrIZV7wvsfivkoUIV?usp=sharingModel SD/SSD](https://drive.google.com/drive/folders/1kmCOspV-hgP4K3lZrIZV7wvsfivkoUIV?usp=sharingModel%20SD/SSD).
- [www 3] [https://drive.google.com/drive/folders/1gyz5N4O0oBKmnaU-Fvk9CoNDt1kX2AKy?usp=sharingWizualizacje dynamiczne BSC](https://drive.google.com/drive/folders/1gyz5N4O0oBKmnaU-Fvk9CoNDt1kX2AKy?usp=sharingWizualizacje%20dynamiczne%20BSC).
- [www 4] [https://drive.google.com/file/d/14pkfyApmy4mQrSLXW2zgulSewyqZsytT/view?usp=sharingTabele równań](https://drive.google.com/file/d/14pkfyApmy4mQrSLXW2zgulSewyqZsytT/view?usp=sharingTabele%20r%C3%B3wna%C5%84).
- [www 5] [https://drive.google.com/file/d/1EmexUL2sQfhT9\\_o6dTaN3wa\\_hvTFIbHu/view?usp=sharingDashboard A, \(B\) modelu SD/SSD](https://drive.google.com/file/d/1EmexUL2sQfhT9_o6dTaN3wa_hvTFIbHu/view?usp=sharingDashboard%20A,%20B%20modelu%20SD/SSD).
- [www 6] [https://drive.google.com/file/d/1xUJC9aN-1HErup3WB\\_MKN3DNNQmld2cg/view?usp=sharingSchemat blokowy SSD WINGS](https://drive.google.com/file/d/1xUJC9aN-1HErup3WB_MKN3DNNQmld2cg/view?usp=sharingSchemat%20blokowy%20SSD%20WINGS).
- [www 7] [https://drive.google.com/file/d/1g0JBTSHSWFiWECrmNJDd6tnPtNERVllu/view?usp=sharingRównania modelu SD/SSD – uogólnione](https://drive.google.com/file/d/1g0JBTSHSWFiWECrmNJDd6tnPtNERVllu/view?usp=sharingR%C3%B3wnania%20modelu%20SD/SSD%20-%20uog%C3%B3lnione).
- [www 8] [https://drive.google.com/drive/folders/1kmCOspV-hgP4K3lZrIZV7wvsfivkoUIV?usp=sharingwarianty SD/SSD](https://drive.google.com/drive/folders/1kmCOspV-hgP4K3lZrIZV7wvsfivkoUIV?usp=sharingwarianty%20SD/SSD).
- [www 9] [https://drive.google.com/file/d/1rOsItgq9ohTDvKhZZtQQ\\_7BTztXG9V9c/view?usp=sharing](https://drive.google.com/file/d/1rOsItgq9ohTDvKhZZtQQ_7BTztXG9V9c/view?usp=sharing).



## **Zastosowanie cen online i modeli autoregresyjnych w krótkoterminowych prognozach cen żywności**

*Jarosław Janecki<sup>1</sup>*

### **Wprowadzenie**

Rosnąca popularność zawierania transakcji online, wzrost wartości transakcji przeprowadzanych online oraz możliwość porównywania cen, jakie oferuje internet, stwarzają nowe możliwości monitorowania zmian cen. Informacje o zmianach cen online w trakcie danego miesiąca pozwalają na ocenę dynamiki inflacji praktycznie w czasie rzeczywistym, przed opublikowaniem wskaźnika CPI (ang. Consumer Price Index) za dany miesiąc. W przypadku krótkoterminowych prognoz zmian cen (w ujęciu miesięcznym) ceny online mogą stanowić ważne źródło informacji pozwalające zmniejszyć błąd prognozy.

Wszystkie wymienione powyżej czynniki odnoszące się do cen online składają do podjęcia badań z wykorzystaniem takich danych z polskiego rynku. Niepewność, z jaką mamy do czynienia w przypadku krótkoterminowych prognoz wskaźnika inflacji (prognozy w ramach konsensusów rynkowych a dane publikowane przez GUS), nie jest czynnikiem pożądanym. Ceny online stwarzają możliwość uzyskania szybszej i pełniejszej wiedzy o procesach cenowych. Ceny żywności stanowią najważniejszy komponent wskaźnika CPI w Polsce. Głównym celem niniejszego badania jest sprawdzenie możliwości zastosowania cen żywności online oraz modeli autoregresyjnych do krótkoterminowych prognoz cen żywności.

---

<sup>1</sup> Szkoła Główna Handlowa w Warszawie, Kolegium Zarządzania i Finansów, Katedra Ekonomii Stosowanej, e-mail: jaroslaw.janecki@sgh.waw.pl.

## 2.1. Przegląd literatury – badania cen online

Zainteresowanie cenami online jest efektem między innymi wzrostu udziału sprzedaży online w łącznej sprzedaży. Ceny online są monitorowane przede wszystkim przez urzędy statystyczne, jak również przez większość banków centralnych. Do najbardziej znanych badań dotyczących cen online należy zaliczyć badania cen online w ramach projektów Billion Prices Project (BPP) oraz Digital Price Index (DPI). W ramach projektu BPP jest zbieranych około 15 milionów danych dziennie od ponad 900 sprzedawców detalicznych w 20 krajach [Cavallo, Rigobon, 2016]. Od 2011 roku dane te pozwalają tworzyć indeksy cen o wysokiej częstotliwości, które są następnie oferowane bankom centralnym i klientom z sektora finansowego. Z kolei stworzony przez Adobe Analytics projekt DPI obejmuje dzienne ceny 80% sprzedawców z listy Fortune 500 [Goolsbee, Klenow, 2018]. Co ważne, dane Adobe Analytics dotyczą rzeczywistych cen transakcji oraz ilości zakupionych produktów.

Generalnie można stwierdzić, że ceny online stanowią uzupełnienie w sposobie pozyskiwania informacji o procesach cenowych [Cavallo, 2012; Hillen, 2019]. Ceny online są coraz powszechniej wykorzystywane przez urzędy statystyczne w celu pomiaru inflacji. Wskazują na to opracowania tych instytucji, w tym między innymi w Stanach Zjednoczonych [Horrigan, 2013], Wielkiej Brytanii [Breton i in., 2015], Holandii [Griffioen, de Haan, Willenborg, 2014], Norwegii [Nygaard, 2015] czy też Nowej Zelandii [Krsinich, 2015]. Należy podkreślić, że wzrost zainteresowania cenami online przez urzędy statystyczne wystąpił szczególnie w okresie pandemii COVID-19. Ceny online stanowiły zazwyczaj dopełnienie badań cen, tymczasem ograniczenie działalności gospodarczej i w niektórych miesiącach brak możliwości przeprowadzenia standardowej procedury sprawdzania cen w placówkach handlowych przyczyniły się do zwrócenia uwagi na ceny online. Było to zgodne z rekomendacjami Eurostatu, który wskazywał na tego rodzaju możliwości [GUS, 2020].

W Polsce ceny online oprócz głównego Urzędu Statystycznego [GUS, 2020] są również przedmiotem badań banku centralnego [Macias, Stelmasiak, 2019]. Ceny online są zbierane, a następnie publikowane w ramach syntetycznego wskaźnika cen online przez CASE [Radzikowski, Śmietanka, 2016].

Badania z wykorzystaniem cen online wskazują na wiele ich zalet, jak również wad [Cavallo, Rigobon, 2016]. Do podstawowych zalet analizy cen online należy zaliczyć: możliwość zbierania danych w sposób zdalny, niski koszt jednej obserwacji, wysoką częstotliwość pozyskiwania danych oraz możliwość zwięk-

szenia zbioru szczegółowych danych dotyczących procesów cenowych. Bardzo ważną zaletą jest możliwość analizy informacji praktycznie w czasie rzeczywistym. Może to wspomagać przewidywanie potencjalnych szoków cenowych, jakie pojawią się w oficjalnych indeksach cen konsumpcyjnych. Ceny online pełnią w tym przypadku rolę wspomagającą w bardziej precyzyjnym i szybszym rozpoznaniu szoku, który pojawi się za chwilę we wskaźniku, który powstaje w sposób tradycyjny. Zbiory danych cen online mogą być użytecznym dodatkiem do modeli prognozowania inflacji. Potwierdzają to między innymi badania na danych z USA i Wielkiej Brytanii [Aparicio, Bertolotto, 2016].

Ważną właściwością cen online jest ich praktyczne wykorzystanie w porównywaniu cen. Indeksy zmian cen budowane z wykorzystaniem cen online w zbliżony sposób oddają zmiany oficjalnych wskaźników inflacji urzędów statystycznych. Różnice, jakie w tym przypadku się pojawiają, są oczywiście efektem różnic metodologicznych stosowanych w konstrukcji indeksów. Budowa wskaźników cen online ma na celu nie tyle odzwierciedlenie zmiany oficjalnych wskaźników inflacji, ile monitorowanie w szybszy sposób zmiany cen w handlu internetowym, co może mieć swoje konsekwencje w tradycyjnych zakupach poza internetem. Austan D. Goolsbee i Peter J. Klenow potwierdzili, że wskaźniki inflacji wyznaczone z wykorzystaniem cen online są niższe od wskaźników inflacji powstałych w tradycyjny sposób w ramach pomiaru cen w sklepach stacjonarnych. Różnica ta wynosi punkt procentowy rocznie w latach 2014-2017 i trzy punkty procentowe rocznie w przypadku uwzględnienia nowych kategorii produktów [Goolsbee, Klenow, 2018]. Różnice te w oczywisty sposób są ważne z punktu widzenia analizy procesów cenowych, jak również mogą mieć konsekwencje dla zmian ogólnego poziomu cen w gospodarce.

Podstawową wadą danych pozyskiwanych w ramach badań cen online są problemy związane z analityką dużych zbiorów danych. Gromadzenie i przetwarzanie danych w postaci dużych zbiorów może się przyczyniać do problemów o charakterze technicznym. Do wad zalicza się również problem określania wag dla poszczególnych kategorii produktów. Ustalenie reprezentantów dla poszczególnych cen jest utrudnione. Zbierane ceny online w większości przypadków są pozbawione informacji o ilości sprzedawanych produktów, co generuje problem związany z prawidłowym uwzględnianiem ich w wydatkach konsumenckich (wyjątek stanowią ceny online pozyskiwane w ramach projektu DPI).

## 2.2. Badania empiryczne

W badaniu wykorzystano dane o cenach żywności online pozyskiwane w ramach projektu CASE w interwałach tygodniowych w okresie od stycznia 2016 roku do lutego 2021 roku (łącznie 64 miesiące). Łącznie w badanym okresie zebrano ponad 1,5 mln cen, średnio ponad 50 tys. reprezentantów cen online w interwale tygodniowym. Do celów badania utworzono 5 szeregów danych: dane z pierwszego tygodnia w danym miesiącu do pierwszego tygodnia z miesiąca poprzedniego (iCPIF1), dane z drugiego tygodnia (iCPIF2), trzeciego tygodnia (iCPIF3), czwartego tygodnia (iCPIF4) oraz miesięczna średnia cen z czterech pomiarów w danym miesiącu (iCPIFa).

W badaniu zastosowano model autoregresyjny z rozkładem opóźnień ARDL (ang. Autoregressive Distributed Lags). Należy on do modeli autoregresyjnych, w których zmienna objaśniana zależy zarówno od zmiennych objaśniających, jak i od własnych opóźnionych wartości. Model ten został szczegółowo przedstawiony w pracach: Pesaran i Smith [1998] oraz Pesaran, Shin i Smith [2001]. Popularność modeli ARDL wynika z możliwości jednoczesnego szacowania długo- i krótkookresowego parametrów modelu, unikając jednocześnie problemów, jakie stwarzają dane niestacjonarne. Model ARDL w kontekście badań z wykorzystaniem zmian cen online został zastosowany między innymi przez Cavallo i Rigobona [2016]. Przeprowadzenie badań z użyciem modelu ARDL oferuje pakiet ekonometryczny Eviews. Program ten został wykorzystany w celu przeprowadzenia obliczeń.

Ogólną postać modelu ARDL przedstawiono wzorem (1):

$$\Delta Y_t = \beta_0 + \sum_{i=1}^p \lambda_i \Delta Y_{t-i} + \sum_{i=0}^q \delta_i \Delta X_{t-i} + \varphi_1 Y_{t-1} + \varphi_2 X_{t-1} + \mu_t \quad (1)$$

gdzie:

- $\Delta Y_t$  – przyrost nielosowej, nieopóźnionej zmiennej endogenicznej w okresie  $t$ ,
- $Y_{t-i}$  – nielosowa, opóźniona zmienna endogeniczna w okresie  $t - i$ ,
- $\Delta X_t$  – przyrost nielosowej, nieopóźnionej zmiennej egzogenicznej w okresie  $t$ ,
- $\Delta X_{t-i}$  – przyrost nielosowej, opóźnionej zmiennej egzogenicznej w okresie  $t - i$ ,
- $\lambda_i$  i  $\delta_i$  – parametry relacji krótkookresowej,
- $\varphi_1$  i  $\varphi_2$  – parametry relacji długookresowej,
- $p, q$  – optymalne opóźnienia,
- $\beta_0$  – stała,
- $\mu_t$  – składnik zakłócający modelu,
- $i = 1, \dots, k$ .

Model ARDL ( $p, q$ ) składa się z dwóch komponentów: krótkookresowego (ST) i długookresowego (LT). Komponent krótkookresowy  $ST$  opisuje równanie (2):

$$ST = \sum_{i=1}^p \lambda_i \Delta Y_{t-i} + \sum_{i=0}^q \delta_i \Delta X_{t-i} \quad (2)$$

Parametry  $\lambda$  i  $\delta$  są niezerowe, ich wartości reprezentują siłę zależności w krótkim okresie. Komponent długookresowy  $LT$  w modelu ARDL składa się z rezydualnych zmiennych opóźnionych  $Y_{t-1}$  oraz  $X_{t-1}$ . Opisuje go równanie (3):

$$LT = \varphi_1 Y_{t-1} + \varphi_2 X_{t-1} \quad (3)$$

Współczynniki  $\varphi_1$  i  $\varphi_2$  określają szybkość dostosowywania do długoterminowej równowagi.

Model ARDL ( $p, q$ ), gdzie  $p$  i  $q$  wskazują na rząd wielomianów operatorów użytych do zapisu modelu. Zauważmy, że dla  $p = 0$  i  $q = 0$  otrzymujemy klasyczny model regresji liniowej.

Przeprowadzono dwie estymacje modelu ARDL (ARDL1 i ARDL2). W ramach pierwszej estymacji (model ARDL1) zmiennymi objaśniającymi były zmiany cen online żywności w poszczególnych czterech tygodniach danego miesiąca (równanie 4):

$$\begin{aligned} \Delta CPIF_t = \beta_0 + \sum_{i=1}^p \lambda_i \Delta CPIF_{t-i} + \sum_{i=0}^{q_1} \delta_i \Delta iCPIF1_{t-i} + \sum_{i=0}^{q_2} \delta_i \Delta iCPIF2_{t-i} + \\ + \sum_{i=0}^{q_3} \delta_i \Delta iCPIF3_{t-i} + \sum_{i=0}^{q_4} \delta_i \Delta iCPIF4_{t-i} + \varphi_1 iCPIF_{t-1} + \\ + \varphi_2 iCPIF1_{t-1} + \varphi_3 iCPIF2_{t-1} + \varphi_4 iCPIF3_{t-1} + \varphi_5 iCPIF4_{t-1} + \mu_t \end{aligned} \quad (4)$$

gdzie:

$\Delta CPIF_t$  – zmiana cen żywności w miesiącu  $t$  (dane GUS),

$\Delta CPIF_{t-1}$  – zmiana cen żywności w miesiącu  $t - 1$ ,

$iCPIF1$ ;  $iCPIF2$ ;  $iCPIF3$ ;  $iCPIF4$  – zmiana internetowych cen żywności odpowiednio w pierwszym, drugim, trzecim i czwartym tygodniu danego miesiąca do internetowych cen żywności z pierwszego tygodnia z miesiąca poprzedniego,

$\beta_0$  – stała,

$p, q_1, q_2, q_3, q_4$  – optymalne opóźnienia,

$\lambda_i$  i  $\delta_i$  – niezerowe parametry krótkookresowej zależności,  
 $\varphi_1, \varphi_2, \varphi_3, \varphi_4, \varphi_5$  – niezerowe parametry długookresowej zależności,  
 $\mu_t$  – składnik zakłócający modelu,  
 $i = 1, \dots, k$ .

W ramach drugiej estymacji (model ARDL2) została uwzględniona miesięczna średnia cen z czterech pomiarów w danym miesiącu (równanie 5):

$$\Delta CPIF_t = \beta_0 + \sum_{i=1}^p \lambda_i \Delta CPIF_{t-i} + \sum_{i=0}^q \delta_i \Delta iCPIFa + \varphi iCPIFa + \mu_t \quad (5)$$

gdzie:

$CPIF_t$  – zmiana cen żywności w danym miesiącu  $t$  (dane GUS),  
 $iCPIFa$  – średnia zmiana internetowych cen żywności miesiąc do miesiąca,  
 $\lambda_i$  i  $\delta_i$  – niezerowe parametry krótkookresowej zależności,  
 $p, q$  – optymalne opóźnienia,  
 $\varphi$  – niezerowy parametr długookresowej zależności,  
 $\beta_0$  – stała,  
 $\mu_t$  – składnik zakłócający modelu.

Badanie stacjonarności za pomocą rozszerzonego testu pierwiastka jednostkowego Dickeya-Fullera (ADF) wskazało, że zmienne są szeregami zintegrowanymi stopnia pierwszego I(1) lub stopnia zerowego I(0). W celu zbadania kointegracji zmiennych przeprowadzono analizę modelu ARDL (bound testing). Testowanie modelu ARDL jest wrażliwe na rząd opóźnienia zmiennych. Korzystając z kryterium informacyjnego Akaike'a (AIC), przyjęto rząd opóźnienia 1. Sformułowano hipotezę  $H_0$  mówiącą o braku związku między zmiennymi. Zatem brak podstaw do odrzucenia hipotezy zerowej występuje w sytuacji, gdy zarówno statystyka  $F$ , jak i statystyka  $t$  są bliższe zera niż wartości krytyczne dla zmiennych I(1). Przeprowadzone obliczenia wskazały, że statystyki  $F$  i  $t$  są wyższe niż wartości krytyczne dla zmiennych I(1) przy poziomie istotności 10%, 5%, 2,5% oraz 1%. Można zatem odrzucić hipotezę zerową o braku związku między zmiennymi. W przypadku obu modeli ARDL1 i ARDL2 stwierdzono występowanie długookresowych relacji pomiędzy internetowymi cenami żywności. Wyniki przeprowadzonych obliczeń dla obu modeli przedstawiono w tabeli 1.

**Tabela 1.** Modele ARDL1 i ARDL2 – (F-bounds test)

| Model ARDL1                                                   |           |       |       | Model ARDL2                |           |       |       |
|---------------------------------------------------------------|-----------|-------|-------|----------------------------|-----------|-------|-------|
| Test statystyczny                                             | Istotność | I(0)  | I(1)  | Test statystyczny          | Istotność | I(0)  | I(1)  |
| F-statystyka: 10,2<br>k: 4                                    | 10%       | 2,45  | 3,52  | F-statystyka: 24,3<br>k: 1 | 10%       | 4,04  | 4,78  |
|                                                               | 5%        | 2,86  | 4,01  |                            | 5%        | 4,94  | 5,73  |
|                                                               | 2,5%      | 3,25  | 4,49  |                            | 2,5%      | 5,77  | 6,68  |
|                                                               | 1%        | 3,74  | 5,06  |                            | 1%        | 6,84  | 7,84  |
| Wielkość próbki: n = 65                                       |           |       |       | Wielkość próbki: n = 65    |           |       |       |
|                                                               | 10%       | 2,57  | 3,68  |                            | 10%       | 4,18  | 4,93  |
|                                                               | 5%        | 3,07  | 4,27  |                            | 5%        | 5,13  | 5,98  |
|                                                               | 1%        | 4,19  | 5,69  |                            | 1%        | 7,32  | 8,44  |
| Wielkość próbki: n = 60                                       |           |       |       | Wielkość próbki: n = 60    |           |       |       |
|                                                               | 10%       | 2,57  | 3,71  |                            | 10%       | 4,15  | 4,95  |
|                                                               | 5%        | 3,06  | 4,31  |                            | 5%        | 5,13  | 6,00  |
|                                                               | 1%        | 4,18  | 5,68  |                            | 1%        | 7,40  | 8,51  |
| Test granic (t-Bounds Test): Hipoteza zerowa: brak zależności |           |       |       |                            |           |       |       |
| t-statystyka: -5,78                                           | 10%       | -2,57 | -3,66 | t-statystyka: -6,69        | 10%       | -2,57 | -2,91 |
|                                                               | 5%        | -2,86 | -3,99 |                            | 5%        | -2,86 | -3,22 |
|                                                               | 2,5%      | -3,13 | -4,26 |                            | 2,5%      | -3,13 | -3,5  |
|                                                               | 1%        | -3,43 | -4,6  |                            | 1%        | -3,43 | -3,82 |

Źródło: Obliczenia własne.

W kolejnym etapie badań sprawdzono wartość błędów prognozy generowaną przez oba modele. Dla obu modeli zostały obliczone trzy rodzaje błędów: błąd średniokwadratowy (RMSE), który jest miarą różnic między wartościami przewidywanymi przez model lub estymator a wartościami zaobserwowanymi, średni błąd bezwzględny (MAE), który jest miarą błędów pomiędzy parami obserwacji wyrażających to samo zjawisko, symetryczny średni bezwzględny błąd procentowy (SMAPE) będący miarą dokładności opartą na błędach procentowych. W przypadku prognoz cen żywności modelem generującym mniejsze błędy jest model ARDL1 uwzględniający zmiany cen online żywności w poszczególnych tygodniach. Wszystkie trzy rodzaje błędów są w tym przypadku mniejsze niż dla modelu ARDL2 (tabela 2).

**Tabela 2.** Błędy prognozy generowane przez modele ARDL1 i ARDL2

| Błędy prognozy                                                                                   | ARDL1    | ARDL2    |
|--------------------------------------------------------------------------------------------------|----------|----------|
| Błąd średniokwadratowy (Root Mean Squared Error, RMSE)                                           | 0,480418 | 0,491191 |
| Średni błąd bezwzględny (Mean Absolute Error, MAE)                                               | 0,382136 | 0,396298 |
| Symetryczny średni bezwzględny błąd procentowy (Symmetric Mean Absolute Percentage Error, SMAPE) | 0,381006 | 0,395098 |

Źródło: Obliczenia własne.

Przeprowadzone obliczenia wskazują na możliwość wykorzystania modeli ARDL w krótkoterminowym modelowaniu cen żywności online. Uzyskiwane błędy prognozy wynikają między innymi ze sposobu ich uwzględniania w modelach.

## Podsumowanie

Ceny online stanowią niezwykle ważny punkt odniesienia w ocenie procesów inflacyjnych zarówno przez producentów, jak i konsumentów. Dają one możliwość zwiększenia naszej wiedzy na temat procesów cenowych zachodzących w gospodarce. Wśród wielu zalet do najważniejszych należy fakt szybkiego pozyskania informacji na temat zmian cen online, które zachodzą w gospodarce. Może to być cenna wiedza zarówno dla decydentów (szczególnie w przypadku realizacji polityki pieniężnej opartej na celu inflacyjnym), jak i samych konsumentów. Na szczególną uwagę zasługuje możliwość konstruowania wskaźników inflacji opartych na cenach online. Przeprowadzone badania wykazały, że do krótkoterminowych prognoz zmian cen żywności w gospodarce można wykorzystać ceny online żywności i modele autoregresyjne ARDL. W kolejnych krokach badanie może zostać rozszerzone na inne kategorie cen. Ich uwzględnienie może się przyczynić do ograniczenia błędu prognozy w krótkoterminowych modelach inflacji i zmniejszenia niepewności, które jest generowane w momencie publikowania oficjalnych danych o inflacji.

## Literatura

- Aparicio D., Bertolotto M. (2016), *Forecasting Inflation with Online Prices*, MIT.
- Breton R., Gareth C., Metcalfe L., Milliken N., Payne C., Winton J., Woods A. (2015), *Research Indexes Using Web Scraped Data*, Office for National Statistics, UK.
- Cavallo A. (2012), *The Billion Prices Project: Building Economic Indicators from Online Data*, MIT Sloan, Geneva, May 31.
- Cavallo A., Rigobon R. (2016), *The Billion Prices Project: Using Online Prices for Measurement and Research*, "The Journal of Economic Perspectives", Vol. 30(2), s. 151-178.
- Goolsbee A.D., Klenow P.J. (2018), *Internet Rising, Prices Falling: Measuring Inflation in a World of E-Commerce*, "AEA Papers and Proceedings", Vol. 108, s. 488-492.
- Griffioen R., de Haan J., Willenborg L. (2014), *Collecting Clothing Data from the Internet*, Proceedings of Meeting of the Group of Experts on Consumer Price Indexes, May 26-28, UNECE.

- 
- GUS (2020), *Wytyczne dotyczące opracowania HICP w kontekście kryzysu związanego z COVID-19*, Nota metodologiczna GUS, 14.04.2020 r.
- Hillen J. (2019), *Web Scraping for Food Price Research*, "British Food Journal", Vol. 121(12), November, s. 3350-3361.
- Horrigan M.W. (2013), *Big Data: A Perspective from the BLS*, "Amstat News", January 1, s. 25-27.
- Krsinich F. (2015), *Price Indexes from Online Data Using the Fixed-Effects Window-Splice (FEWS) Index*, Paper presented at the Ottawa Group, Tokyo, Japan, May 20-22.
- Macias P., Stelmasiak D. (2019), *Food Inflation Nowcasting with Web Scraped Data*, NBP Working Paper No. 302, Warsaw.
- Nygaard R. (2015), *The Use of Online Prices in the Norwegian Consumer Price Index*, Paper prepared for the meeting of the Ottawa Group, Tokyo, Japan, May 20-22.
- Pesaran M.H., Shin Y., Smith R.J. (2001), *Bounds Testing Approaches to the Analysis of Level Relationships*, "Journal of Applied Econometrics", Vol. 16, Iss. 3, s. 289-326.
- Pesaran M.H., Smith R.P. (1998), *Structural Analysis of Cointegrating VARs*, "Journal of Economic Surveys", Vol. 12, Iss. 5, s. 471-505.
- Radzikowski B., Śmietanka A. (2016), *Online CASE CPI*, First International Conference on Advanced Research Methods and Analytics, CARMA 2016, Universitat Politècnica de Valencia, Valencia.

## Teoria ujawnionych preferencji i wybory międzyokresowe w warunkach ryzyka

*Józef Stawicki<sup>1</sup>*

### Wprowadzenie

Pionierem teorii ujawnionych preferencji<sup>2</sup> był amerykański ekonomista Paul Samuelson<sup>3</sup>. Teoria ta opiera się na założeniu, że preferencje konsumentów ujawniają się w ich decyzjach zakupowych. Dotychczasowe badania teoretyczne szły w kierunku od preferencji do decyzji. Naukowcy działający w tym obszarze wiedzy opracowali złożone i wyrafinowane modele matematyczne, aby uchwycić preferencje, które są „ujawniane” poprzez zachowania konsumentów w zakresie wyborów konsumpcyjnych. Problem dotyczy nie tylko obszaru konsumpcji. Teoria podejmowania decyzji stoi przed podobnym dylematem. Jeśli znamy preferencje decydenta, to za pomocą opracowanych metod można wskazać właściwe decyzje jako rozwiązanie zagadnienia jedno- czy wielokryterialnego. Ten kierunek badań w zakresie wielokryterialnych decyzji jest bardzo rozwinięty [Trzaskalik, 2014].

Pytanie, jakie trzeba sobie postawić, to czy na podstawie obserwowanych i podejmowanych decyzji można określić preferencje decydenta (jego funkcję użyteczności). Zakłada się przy tym racjonalność decydenta (konsumenta). Postawiony problem został dobrze opisany i przedstawionych jest wiele rozwiązań. Doskonałym przeglądem dotychczasowego dorobku w tym zakresie, stanowiącym niejako podsumowanie, jest monografia C.P. Chambersa i F. Echenique’a [2016]. Stan badań w zakresie klasycznego rozumienia teorii preferencji ujawnionych oraz możliwe kierunki dalszych badań przedstawiono między innymi w pracy Echenique’a [2019].

---

<sup>1</sup> Uniwersytet Mikołaja Kopernika w Toruniu, Wydział Nauk Ekonomicznych i Zarządzania, Katedra Zastosowań Informatyki i Matematyki w Ekonomii, e-mail: stawicki@umk.pl.

<sup>2</sup> Teoria jest znana pod anglojęzyczną nazwą „revealed preference theory”. W polskiej literaturze stała się powszechna dzięki podręcznikowi Variana [1995].

<sup>3</sup> Pierwsza praca na ten temat stała się podstawą bardzo szerokiej dyskusji. Zob. Samuelson [1938, s. 61-71].

Teoria ujawnionych preferencji, choć nie ma długiej historii w ekonomii, to jednak odgrywa istotną rolę, nadając empiryczne znaczenie hipotezie racjonalności w ekonomii neoklasycznej. Badania rozpoczęte przez Samuelsona [1938] i kontynuowane przez Afriata [1967], Diewerta [1973] i Variana [1982] wyjaśniły behawioralne znaczenie hipotezy, że konsument maksymalizuje użyteczność. Można powiedzieć, że teoria ta poprzedziła teorię finansów behawioralnych i stała się początkiem ekonomii eksperymentalnej. Teoria osiąga punkt kulminacyjny w słynnym twierdzeniu Afriata zwanym „Słabym Aksjomatem Preferencji Ujawnionych – WARP”, które stwierdza, że empiryczna treść racjonalnego zachowania (maksymalizacja użyteczności przez lokalnie nienasyconego konsumenta) jest taka sama, jak w przypadku konsumenta z jednostajną rosnącą i wklęsłą funkcją użyteczności. Z kolei zbiory danych, które są zgodne z teorią, to te, które spełniają „Ogólny Aksjomat Preferencji Ujawnionych – GARP” Afriata (GARP; terminologia pochodzi od Variana [1982]). Analizując teorię preferencji ujawnionych, warto zwrócić uwagę na inne prace pozostające może poza głównym nurtem badań. Do nich należą niewątpliwie prace Kennetha Arrowa [Arrow, 1951; Arrow, 1959]. Wiele problemów badanych współcześnie dotyczy wyborów na gruncie społecznym. Drugim kierunkiem mało docenionym w literaturze naukowej przedmiotu są prace Ludwiga von Auera [von Auer, 1998; von Auer, 2004]. Wyniki tych prac można powiązać z teorią programowania dynamicznego. Na koniec trzeba wspomnieć o współczesnych osiągnięciach Matthew Polissona i jego współpracowników [Polisson, Quah, Renou, 2020]. Kierunek badań teorii preferencji ujawnionych w warunkach ryzyka i niepewności dotyczy współczesnych problemów i jest najnowszym kierunkiem badawczym. Należy jeszcze wspomnieć o wyborach międzyokresowych ściśle związanych z dynamicznym podejściem we wspomnianej teorii. Wybór międzyokresowy należy rozważać jako rozszerzenie zastosowania teorii oczekiwanej użyteczności i zdyskontowanie jej na jeden okres z uwzględnieniem wielu towarów lub stanów. Niewątpliwie osiągnięcia w tym zakresie należą do Martina Browninga [1989]. Wiele ciekawych propozycji i rozwiązań zaproponowano również w pracy von Auera [2004]. Problem staje się skomplikowany, gdy weźmie się pod uwagę sprawiedliwość społeczną międzypokoleniową i wybory decydujące o rozwoju w długim okresie. Niektóre spostrzeżenia związane z tą konfrontacją przedstawił Rybicki [2014]. W pracy tej znalazło się odniesienie do wartości dyskontowania „dalekiej przyszłości” według formuł wykładniczych oraz hiperbolicznych, znanych z pracy Harveya [1995].

### 3.1. Interdyscyplinarność teorii preferencji ujawnionych

Teoria preferencji ujawnionych jest problemem interdyscyplinarnym. Sama teoria powstała na gruncie ekonomicznym [patrz: Samuelson, 1938]. Problemy dotyczyły jednak przynajmniej czterech dyscyplin:

- 1) ekonomii,
- 2) psychologii,
- 3) filozofii i metodologii nauk,
- 4) matematyki.

Wynika to z faktu wielokierunkowego powiązania ekonomii z innymi dyscyplinami, choć niektóre z twierdzeń dotyczą w specyficzny sposób na przykład tylko psychologii.

Główny problem ekonomiczny dotyczył warstwy poznawczej zachowań konsumpcyjnych. Jednak w warstwie praktycznej najistotniejszym aspektem jest prognozowanie popytu i dostosowanie rynku do potrzeb konsumenta. Problem decyzji pojedynczej osoby przekłada się na problem decyzji społecznych i wyborów społecznych bez względu na to, czy jest to decyzja zbiorowa, czy indywidualna podejmowana w imieniu zbiorowości. Dotyczy to także decyzji innych niż konsumpcyjne. Szczególnie ważne są tu decyzje produkcyjne, decyzje organizacyjne, decyzje inwestycyjne czy decyzje dotyczące rozwoju infrastruktury<sup>4</sup>. Zachowania konsumentów mają istotne znaczenie dla prognozowania popytu i analizy rynku. W tym miejscu warto wspomnieć o wielu wysiłkach i skutecznych metodach w zakresie łączenia danych z ujawnionych preferencji i deklarowanych preferencji [patrz: Rybicka, 2015; Bąk, 2004; Bartołomowicz, 2009]. Zagadnienie jest bardzo praktyczne i znanych jest wiele empirycznych prac na ten temat [Birol, Kontoleon, Smale, 2006; Paccagnan, 2007].

Odrębnym obszarem, który należałoby powiązać z zagadnieniem preferencji ujawnionych, są niewątpliwie finanse behawioralne z fundamentalną teorią perspektywy i wszystkimi heurystykami przedstawianymi w tym przedmiocie. Jest to w zasadzie domeną psychologii ze specyficznym podejściem w ekonomii i zarządzaniu<sup>5</sup>. Inną dziedziną, która jest dziś mocno rozwijana, jest ekonomia eksperymentalna. Na styku eksperymentu i preferencji ujawnionych rodzi się wiele pytań<sup>6</sup>.

<sup>4</sup> Przegląd metod w tym obszarze jest prezentowany w wielu monografiach Katowickiej Szkoły Podejmowania Decyzji stworzonej i skutecznie rozwijanej przez prof. Tadeusza Trzaskalika. Wśród nich można choćby wymienić takie monografie, jak: Trzaskalik, red. [2014]; Trzaskalik [1990] czy Trzaskalik [1986].

<sup>5</sup> Warto tu wspomnieć o wielu pracach polskich uczonych: Tyszka [2010]; Tyszka, red. [2004]; Zielonka [2008].

<sup>6</sup> Świadomi tych pytań są niewątpliwie prof. Ewa Roszkowska i prof. Tomasz Wachowicz oraz ich współpracownicy – Autorzy wielu monografii i artykułów z zakresu wspomaganego nego-

Wielu autorów [patrz np. Boland, 2001] zadawało sobie pytanie „czy jest możliwa falsyfikacja teorii preferencji ujawnionych?”. Niemożliwość falsyfikacji czyni w jakimś sensie bezużyteczne rozważania na temat preferencji ujawnionych. Już wcześniej pojawiła się dyskusja na temat samego pojęcia „preferencje”. Daniel Hausman [2000] argumentuje, że to, co ekonomiści nazywają preferencją, jest technicznym pojęciem (nazywa je „preferencją\*”). Pojęcie to różni się od rzeczywistego (psychologicznego znaczenia) preferencji. Według Hausmana *preferencja\** pociąga za sobą zarówno *preferencję*, jak i *przekonanie*. Dyskusja z Hausmanem dotyczyła filozofii nauki. Uznając ważność i jednocześnie atrakcyjność zapożyczoną z teorii modeli matematycznych metody aksjomatycznej, zwraca się jednocześnie uwagę na oderwanie teorii od rzeczywistych problemów ekonomicznych. W tym kontekście należy także spytać o daleko idący formalizm matematyczny [Hands, 2017]. Zdanie „Książka jest napisana w formalnym stylu, a wszystkie główne wyniki są przedstawione jako twierdzenia z dokładnymi dowodami” jest pochwałą, ale być może wiąże się to z wypowiedzią Wassily Leontiewa, który w 1973 otrzymał Nagrodę Banku Szwecji im. Alfreda Nobla w dziedzinie ekonomii. W swoim inauguracyjnym przemówieniu obejmowania funkcji Prezesa Amerykańskiego Stowarzyszenia Ekonomicznego powiedział, że „Bezkrytyczny entuzjazm dla sformułowań matematycznych sprawia niejednokrotnie, że efemeryczność istotnej treści wywodu zostaje ukryta za przerażającą fasadą znaków algebraicznych...”<sup>7</sup>.

Do ujawnionych preferencji potrzeba dynamicznego podejścia. Słaby aksjomat ujawnionych preferencji nie wystarcza do określenia doskonałej zgodności między obserwowanymi wyborami a nieobserwowalnymi ujawnionymi preferencjami. Niezbędne i konieczne jest określenie warunków koniecznych i wystarczających dla tej zgodności. Dynamiczne wybory to wybory międzyokresowe, decyzje, w których rozłożenie w czasie kosztów i korzyści zmienia się w czasie. Szczególne pytanie stawiane nie tylko dziś, to na przykład ile odkładać na emeryturę czy jak inwestować. Problem inwestycji wiąże się z programowaniem dynamicznym – znaną i skutecznie stosowaną metodą wsparcia w podejmowaniu decyzji. Można wskazać wiele innych ciągów decyzji w czasie. Takim przykładem może być wybór ścieżki edukacyjnej<sup>8</sup>. Te pytania mają mocne podstawy w teorii wyborów międzyokresowych.

---

cjacji na wszystkich jego etapach. Jednym z elementów prowadzonych przez nich badań są wnioski z eksperymentów negocjacyjnych [patrz: Roszkowska, Wachowicz, red., 2016; Roszkowska, Filipowicz-Chomko, Wachowicz, 2018; Stawicki i in., 2020]. Sama ekonomia eksperymentalna stała się odrębnym przedmiotem wykładanym na kierunkach ekonomicznych.

<sup>7</sup> Słowa te przywołuje w swojej monografii *Prawda kontra precyzja w ekonomii* [Mayer, 1996, s. 13].

<sup>8</sup> Procesy takie były przedmiotem zainteresowania zespołu prof. Urszuli Sztanderskiej [patrz: Grotkowska, Sztanderska, red., 2015].

W związku z tak postawionym zagadnieniem rodzi się bardzo wiele pytań. Poniżej podano tylko niektóre z nich:

Jak zmienia się funkcja użyteczności w kolejnych krokach?

Czy relacja, jaką posługuje się agent, spełnia warunki relacji preferencji?

Czy bez wsparcia analityka agent jest w stanie podjąć „właściwą” decyzję?

Jeśli podejmuje decyzje, to jak odbiega ona od racjonalnej?

Czy wybory międzyokresowe to to samo, co programowanie dynamiczne?

Czy dokonywane wybory rzeczywiście ujawniają preferencje?

Czy eksperyment ekonomiczny odkrywa preferencje, jakimi posługuje się decydent (agent)?

Co dzieje się w przypadku, gdy nie działa postulat niedosytu nieodzowny w teorii konsumenta?

Co w przypadku, gdy cecha, ze względu na którą buduje się preferencje, nie ma charakteru stymulanty czy destymulanty, a nominanty?

Podstawowym problemem staje się badanie zgodności preferencji deklarowanych z preferencjami ujawnionymi. Preferencje deklarowane mogą być skutecznie badane metodą conjoint<sup>9</sup> (ang. conjoint analysis, CA). Wielu ekonomistów wskazuje na dwie najpopularniejsze w ekonomii metody badania preferencji deklarowanych: metodę wyceny warunkowej (ang. contingent valuation method, CVM) oraz metodę wyboru warunkowego (ang. discrete choice experiment, DCE)<sup>10</sup>.

Zasadniczym przedmiotem badań empirycznych (czy raczej wskazywania ilustracji problemu) są wybory konsumenckie. Uwzględniane są określone koszyki dóbr. Temat tak określony mieścił się w nauczaniu mikroekonomii. Istotniejsze są, jak już wspomniano wyżej, zagadnienia podejmowania decyzji. Badania w tym zakresie są jeszcze na początku. Dotyczą najczęściej decyzji wyboru mieszkania, decyzji wyboru kierunku i miejsca studiów, decyzji wyboru procedury leczenia.

---

<sup>9</sup> Metodę tę spopularyzowali w polskiej literaturze i rozwinęli Walesiak [1996]; Walesiak, Bąk [1998]; Strojny [1998].

<sup>10</sup> Próbę konfrontacji dokonanego wyboru oraz preferowanego wyboru podjął autor niniejszego rozdziału na zajęciach dydaktycznych w 2018 roku w grupie studentów w ramach przedmiotu Analiza preferencji i ryzyka. Studenci po zapoznaniu się z metodologią conjoint analysis zaproponowali badanie wyboru miejsca studiów. Brak reprezentatywności i skąpy materiał statystyczny nie pozwoliły na przygotowanie publikacji na ten temat. Konfrontacja ta wskazała na znaczne rozbieżności preferowanego wyboru oraz dokonanego wcześniej wyboru miejsca studiów (wybór WNEiZ UMK), co było zaskoczeniem dla grupy studentów opracowujących przedstawiony przykład.

### 3.2. Doskonała zgodność między wyborami i preferencjami

Aby zrozumieć problem teorii preferencji ujawnionych, należy zadać sobie zasadnicze pytanie o „doskonałą zgodność między wyborami i preferencjami”. Ludwig von Auer [1998; 2004] przedstawił to w postaci dwóch postulatów. Odnoszą się one do sytuacji deterministycznej bez wprowadzania losowości i niepewności.

POSTULAT 1. Załóżmy, że  $R$  jest „autentyczną (prawdziwą) relacją preferencji” agenta, która generuje obserwowalne wybory  $C$ . Jeśli skonstruowano ujawnione preferencje  $\bar{R}$  na podstawie  $C$ , to powinno to generować  $R = \bar{R}$  ( $C$  – obserwowany wybór,  $R$  – relacja preferencji,  $\bar{R}$  – ujawnione preferencje: jeśli  $C \rightarrow \bar{R}$ , to  $R = \bar{R}$ ).

Postulat 1, rozpatrywany oddzielnie, mówi tylko o części tej idealnej zgodności.

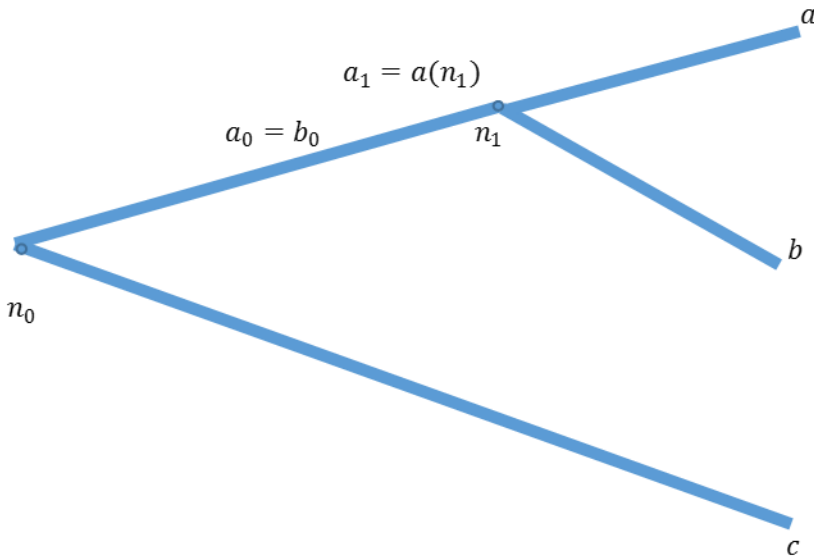
POSTULAT 2. Załóżmy, że  $C$  opisuje obserwowane wybory agenta i ujawnione preferencje  $\bar{R}$  są zbudowane z  $C$ . Jeżeli wygenerowany bazowy wybór  $\bar{C}$  opiera się na ujawnionych preferencjach  $\bar{R}$ , to powinniśmy otrzymać, że  $C = \bar{C}$ .

Tylko wtedy, gdy oba te postulaty zostaną spełnione, można się swobodnie przełączać między pojęciami preferencji i wyboru. Jeśli ta doskonała zgodność między wyborami a preferencjami nie może zostać ustalona, wtedy idea nieobserwowalności preferencji ujawnianych przez obserwowalne wybory traci swój sens.

Pytanie, jak określić zgodność preferencji i wyborów w sytuacji ryzyka i niepewności, pozostaje otwarte. Oczywiście przejście od preferencji do decyzji jest prostym zadaniem rozważanym choćby na zajęciach z mikroekonomii czy ekonomii matematycznej<sup>11</sup>. Trudniejsze jest określenie w kierunku przeciwnym. Ludwig von Auer wykorzystuje przedstawione w literaturze propozycje Richtera, Uzawy oraz Arrowa i dostosowuje je do określenia relacji w przypadku dynamicznym. Analiza jest prowadzona w warunkach pewności i pełnej informacji dla drzewa decyzyjnego. Drzewa decyzyjne składające się z gałęzi, węzłów decyzyjnych i węzłów końcowych są używane do reprezentowania sekwencyjnych wyborów dokonywanych przez agenta w dyskretnych punktach w czasie  $t = 0, 1, 2, \dots, T$ . To drzewo decyzyjne wyrasta z początkowego węzła decyzyjnego  $n_0$  i kończy się skończonym zestawem węzłów  $X = (a, b, \dots, z)$ , które reprezentują ostateczne wybory dokonane przez agenta. Przykładem takim może

<sup>11</sup> Należałoby w tym miejscu przywołać doskonale monografie i podręczniki powstałe w Poznaniu pod kierunkiem profesora Emila Panka [2003; 2005].

być proces decyzyjny zobrazowany na rys. 1. Każdy węzeł końcowy jest połączony z węzłem początkowym przez jedną pełną gałąź. Na przykład kompletna gałąź łącząca węzeł końcowy  $a$  z początkowym węzłem jest oznaczona przez  $a(n_0)$ , który można podzielić na  $T + 1$  części  $a_t$ , gdzie każda zaczyna się w innym węźle decyzyjnym  $n_t$ . W  $n_t$  są określone pozostałe podgałęzie, gdzie ciąg  $a(n_t)$  składa się z części  $(a_t, a_{t+1}, \dots, a_T)$ .



**Rys. 1.** Proces decyzyjny

Propozycja Ludwiga von Auera polega na tak zwanym przerzedzaniu drzewa decyzyjnego, zaczynając od ostatniej decyzji. Proces przerzedzania jest realizowany w następujący sposób:

- (1) rozpocząć od węzła końcowego należącego do poddrzewa  $\chi(n_t)$ ,
- (2) wyeliminować fragmenty jego kompletnej gałęzi, przesuując się wstecz w czasie,
- (3) zatrzymać ten proces eliminacji w pierwszym węźle decyzyjnym, do którego dołącza się kolejna gałąź.

Poddrzewa  $\chi(n_t)$  powstałe w wyniku przerzedzania są oznaczone jako  $A(n_t)$ . Wyobrażając sobie agenta znajdującego się w węźle  $n_t$ , można powiedzieć, że stoi on przed zbiorem możliwości  $A(n_t)$ . Wybór zbioru  $C[A(n_t)]$ , podzbioru  $A(n_t)$ , oznacza dla agenta zestaw opcji, które prawdopodobnie użyje ten agent.

Realizując postulat od wyborów do preferencji i wykorzystując propozycje Richtera [1966], relację preferencji można określić następująco:

- (a)  $a\bar{R}(n_t)b$  wtedy i tylko wtedy, gdy dla dowolnego  $A(n_t)$  z  $\{a(n_1), b(n_1)\} \subseteq A(n_t): a(n_t) \in C[A(n_t)]$ ;
- b)  $a\bar{P}(n_t)b$  wtedy i tylko wtedy, gdy  $a\bar{R}(n_t)b$  i  $\sim a\bar{R}(n_t)b$ ;
- c)  $a\bar{I}(n_t)b$  wtedy i tylko wtedy, gdy zarówno  $a\bar{R}(n_t)b$ , jak i  $b\bar{R}(n_t)a$ .

Sekwencyjność decyzji przywołuje postępowanie decyzyjne, jakim jest programowanie dynamiczne. Mamy tutaj do czynienia z formalnym podejściem zbudowania relacji preferencji. Proces ten jest o tyle trudny, że podjętej decyzji cofnąć nie można bądź jest ona obciążona wysokimi kosztami. Jeśli relacja preferencji została zbudowana na podstawie jednego kryterium, odkrycie takiej relacji nie przedstawia zwykle kłopotu. Jednak procesy decyzyjne i dokonywane przez decydenta wybory opierają się na wielu kryteriach. Stąd propozycje wskazywania decydentowi rozwiązań optymalnych i suboptymalnych czy prawie optymalnych, jak określa profesor Tadeusz Trzaskalik [2015]. Dla rozpoznania prawdziwej relacji można dokonać eksperymentu. Nie jest on jednak możliwy czy wręcz niewskazany w wielu przypadkach. Idea interaktywnych metod wielokryterialnych wspomagania decyzji jest ze wszech miar zasadna. Wskazywanie w procesie interaktywnym rozwiązań nawiązuje niejako do eksperymentu. Procedura STEM opisana i rozszerzona przez profesora Macieja Nowaka [2008] ma znamiona poszukiwania właściwej relacji preferencji, którą posługuje się decydent.

Wiele decyzji wieloetapowych jest podejmowanych bez wsparcia (matematycznego). Wspomnianym już przykładem jest wybór studiów pierwszego stopnia, wybór studiów drugiego stopnia czy wybór studiów podyplomowych. Czy w sekwencji tych wyborów (preferencji ujawnionych) można dojrzeć relację preferencji stojącą u podstaw tych wyborów, pozostaje otwartym pytaniem.

Być może teorii preferencji, zarówno tej dotyczącej preferencji deklarowanych, jak i preferencji ujawnionych, należy się przyjrzeć przez pryzmat teorii „działania ludzkiego” zakorzenionej mocno w austriackiej szkole ekonomicznej<sup>12</sup>. Podstawową kategorią w tym rozumowaniu jest preferencja czasowa i oczywiście działanie rozumiane jako chęć i podjęcie czynów dla zmiany istniejącej sytuacji na lepszą w rozumieniu podejmującego działanie. Krytykę teorii preferencji ujawnionych zapoczątkowanej przez Samuelsona poddał Murray N. Rothbard [1956, 2008]. Zaproponowana przez niego teoria „preferencji de-

<sup>12</sup> Sformułowanie nawiązuje do epokowego dzieła Ludwiga von Misesa pt. *Ludzkie działanie. Traktat o ekonomii* [von Mises, 2011].

monstrowanych” wychodziła naprzeciw oczekiwaniom austriackiej szkoły ekonomii i była zgodna z jej głównymi postulatami. Współcześnie teoria ta znajduje zwolenników wśród uczonych preferujących punkt widzenia austriackiej szkoły ekonomii [Block, Barnett, 2012; Wysocki, Megger, 2019].

## Podsumowanie

Celem rozdziału było przedstawienie problemu preferencji ujawnionych jako dynamicznie rozwijającej się teorii. Powiązanie tych zagadnień z innymi ważnymi z punktu widzenia nie tylko ekonomii, ale głównie psychologii jest potrzebą chwili ze względu na rozwój metod eksperymentalnych w ekonomii i rozwój szkoły ekonomii behawioralnej. W rozdziale postawiono wiele pytań, na które nie ma dziś jasnych i prostych odpowiedzi. Jeśli przedstawione treści wywołały niedosyt wiedzy w tym zakresie albo wzbudziły niepokój, to cel rozdziału został osiągnięty.

## Literatura

- Afriat S.N. (1967), *The Construction of Utility Functions from Expenditure Data*, “International Economic Review”, Vol. 8, s. 67-77.
- Arrow K.J. (1951), *Social Choice and Individual Values*, Yale University Press, New Haven, CT.
- Arrow K.J. (1959), *Rational Choice Functions and Orderings*, “Econometrica”, Vol. 26(102), s. 121-127.
- Bąk A. (2004), *Dekompozycyjne metody pomiaru preferencji w badaniach marketingowych*, Wydawnictwo AE, Wrocław.
- Bartłomowicz T. (2009), *Metody prognozowania preferencji ujawnionych i wyrażonych*, „Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu. Ekonometria”, t. 25, nr 65: Zastosowanie metod ilościowych, s. 141-151.
- Birol E., Kontoleon A., Smale M. (2006), *Combining Revealed and Stated Preference Methods to Assess the Private Value of Agrobiodiversity in Hungarian Home Gardens*, IFPRI Division Discussion Papers, International Food Policy Research Institute.
- Block W.E., Barnett W. (2012), *Transitivity and the Money Pump*, “The Quarterly Journal of Austrian Economics”, Vol. 15(2), s. 237-251.
- Boland L. (2001), *On the Futility of Criticizing the Neoclassical Maximization Hypothesis*, “American Economic Review”, Vol. 71(5), s. 1031-1036.

- Browning M. (1989), *A Nonparametric Test of the Life-cycle Rational Expectations Hypothesis*, "International Economic Review", Vol. 30, s. 979-992.
- Chambers C.P., Echenique F. (2016), *Revealed Preference Theory*, Cambridge University Press, Cambridge.
- Diewert W.E. (1973), *Afriat and Revealed Preference Theory*, "The Review of Economic Studies", Vol. 40, s. 419-425.
- Echenique F.E. (2019), *New Developments in Revealed Preference Theory: Decisions under Risk, Uncertainty, and Intertemporal Choice*, "Annu. Rev. Econ.", Vol. 12, s. 1-30.
- Grotkowska G., Sztanderska U., red. (2015), *Spoleczne i ekonomiczne uwarunkowania wyborów osób w wieku 19-30 lat dotyczących studiowania. Raport końcowy*, <http://produkty.ibe.edu.pl/docs/raporty/ibe-ee-raport-spol-i-ekon-uwar-wyb-19-30-metodolog.pdf> (dostęp: 24.05.2021).
- Hands D.W. (2017), *Revealed Preference, Afriat's Theorem, and Falsifiability: A Review Essay on Revealed Preference Theory by Christopher P. Chambers and Federico Echenique*, [https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=3045952](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3045952) (dostęp: 24.05.2021).
- Harvey Ch. (1995), *Proportional Discounting of Future Costs and Benefits*, „Mathematics of Operations Research”, No. 20, s. 381-399.
- Hausman D.M. (2000), *Revealed Preference, Belief, and Game Theory*, "Economics and Philosophy", Vol. 16(01), s. 99-115.
- Mayer T. (1996), *Prawda kontra precyzja w ekonomii*, Wydawnictwo Naukowe PWN, Warszawa.
- Nowak M. (2008), *Interaktywne wielokryterialne wspomaganie decyzji w warunkach ryzyka. Metody i zastosowania*, Wydawnictwo Akademii Ekonomicznej im. Karola Adamieckiego, Katowice.
- Paccagnan V. (2007), *On Combining Stated Preferences and Revealed Preferences Approaches to Evaluate Environmental Resources Having a Recreational Use*, Università Commerciale Luigi Bocconi, Working Paper No. 4, May.
- Panek E. (2003), *Ekonomia matematyczna*, Wydawnictwo Uniwersytetu Ekonomicznego, Poznań.
- Panek E., red. (2005), *Podstawy ekonomii matematycznej*, Wydawnictwo Uniwersytetu Ekonomicznego, Poznań.
- Polisson M., Quah J.K.-H., Renou L. (2020), *Revealed Preferences over Risk and Uncertainty*, "American Economic Review", Vol. 110(6), s. 1-40.
- Richter M.A. (1966), *Revealed Preference Theory*, "Econometrica", Vol. 34, No. 3 (July), s. 635-645.
- Roszkowska E., Filipowicz-Chomko M., Wachowicz T. (2018), *Ocena akceptowalności wybranych metod wielokryterialnych – badanie eksperymentalne*, Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu, nr 507, s. 219-227.

- Roszkowska E., Wachowicz T., red. (2016), *Negocjacje. Analiza i wspomaganie decyzji*, Wolters Kluwer, Warszawa.
- Rothbard M.N. (2008) [1956], *Toward a Reconstruction of Utility and Welfare Economics* [w:] M. Sennholz (ed.), *On Freedom and Free Enterprise: Essays in Honor of Ludwig von Mises*, The Ludwig von Mises Institute, Auburn, s. 224-262.
- Rybicka A. (2015), *Połączenie danych o preferencjach ujawnionych i wyrażonych*, „Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu”, nr 207, s. 262-270.
- Rybicki W. (2014), *Modelowanie preferencji a zagadnienie zrównoważonego (trwałego) rozwoju* [w:] T. Trzaskalik (red.), *Modelowanie preferencji a ryzyko '14*, Wydawnictwo Uniwersytetu Ekonomicznego, Katowice, s. 127-159.
- Samuelson P.A. (1938), *A Note on the Pure Theory of Consumer's Behaviour*, „Econometrica”, Vol. 5, s. 61-71.
- Stawicki J., Gaspars-Wieloch H., Folipowicz-Chomko M., Roszkowska E., Wachowicz T. (2020), *Profil decydenta wobec ryzyka i niepewności*, Wydawnictwo Naukowe Uniwersytetu Mikołaja Kopernika, Toruń.
- Strojny J. (1998), *Badanie preferencji konsumentów artykułów spożywczych na przykładzie jaj* [w:] T. Trzaskalik (red.), *Modelowanie preferencji a ryzyko '98*, Wydawnictwo Akademii Ekonomicznej im. Karola Adamieckiego, Katowice.
- Trzaskalik T., red. (2014), *Wielokryterialne wspomaganie decyzji. Metody i zastosowania*, PWE, Warszawa.
- Trzaskalik T. (1986), *Wybrane problemy programowania dynamicznego*, Akademia Ekonomiczna, Katowice.
- Trzaskalik T. (1990), *Wielokryterialne dyskretne programowanie dynamiczne. Teoria i zastosowania w praktyce gospodarczej*, Akademia Ekonomiczna, Katowice.
- Trzaskalik T. (2014), *Wielokryterialne wspomaganie decyzji. Przegląd metod i zastosowań*, „Zeszyty Naukowe Politechniki Śląskiej, Seria: Organizacja i Zarządzanie”, z. 74, s. 53-62.
- Trzaskalik T. (2015), *Strategie optymalne i prawie optymalne w dyskretnym stochastycznym programowaniu dynamicznym*, „Journal of Management and Finance”, Vol. 13, No. 4(2), s. 287-300.
- Tyszka T. (2010), *Decyzje. Perspektywa psychologiczna i ekonomiczna*, Wydawnictwo Naukowe „Scholar”, Warszawa.
- Tyszka T., red. (2004), *Psychologia ekonomiczna*, Gdańskie Wydawnictwo Psychologiczne, Gdańsk.
- Uzawa H. (1956), *A Note on Preference and Axioms of Choice*, „Annals of the Institute of Statistical Mathematics”, Vol. 8, s. 35-40.
- Varian H.R. (1982), *The Nonparametric Approach to Demand Analysis*, „Econometrica”, Vol. 50, Iss. 4, s. 945-973.
- Varian H.R. (1995), *Mikroekonomia. Kurs średni, ujęcie nowoczesne*, Wydawnictwo Naukowe PWN, Warszawa.

- 
- Von Auer L. (1998), *Dynamic Preferences, Choice Mechanisms and Welfare*, Springer.
- Von Auer L. (2004), *Revealed Preferences in Intertemporal Decision Making*, "Theory and Decision", Vol. 56(3), s. 269-290.
- Von Mises L. (2011), *Ludzkie działanie. Traktat o ekonomii*, Instytut Ludwiga von Misesa, Warszawa.
- Walesiak M. (1996), *Metody analizy danych marketingowych*, PWN, Warszawa.
- Walesiak M., Bąk A. (1998), *Realizacja badań marketingowych metodą conjoint analysis z wykorzystaniem pakietu statystycznego SPSS for Windows*, Wydawnictwo AE im. Oskara Langego, Wrocław.
- Wysocki I., Megger D. (2019), *Austrian Welfare Economics: A Critical Approach*, "Ekonomia – Economic Review", Vol. 25/1, Wrocław, s. 73-80.
- Zielonka P. (2008), *Behawioralne aspekty inwestowania na rynku papierów wartościowych*, CeDeWu, Warszawa.



## **Część II**



# **RYZYO W ZARZĄDZANIU PROJEKTAMI**





## Zarządzanie ryzykiem projektu w małych instytucjach kultury z wykorzystaniem Scruma

*Dorota Kuchta<sup>1</sup>, Alicja Krawczyńska<sup>2</sup>*

### Wprowadzenie

Instytucje kultury w Polsce prowadzą działalność w różnych formach organizacyjnych, takich jak domy i centra kultury, biblioteki, muzea, teatry i inne. Z uwagi na klarowność rozumienia pojęcia „instytucje kultury” należy podkreślić, że w polskim ustawodawstwie stanowiącym o działalności kulturalnej termin ten jest tożsamy z pojęciem „publicznej instytucji kultury”. Zgodnie z tym podejściem w tekście, używając terminu instytucje kultury, odnoszono się wyłącznie do publicznych podmiotów prowadzących działalność kulturalną. Zamiennie stosowano także pojęcie „organizacja” jako szersze w obrębie nauk o zarządzaniu.

Działalność statutowa instytucji kultury jest rozumiana jako tworzenie, upowszechnianie i ochrona kultury; określa ją Ustawa o organizowaniu i prowadzeniu działalności kulturalnej [1991]. W ramach tej działalności, a także poza nią, instytucje kultury, dzięki dodatkowym środkom pozyskiwanym zarówno z programów uruchamianych corocznie przez Ministra Kultury i Dziedzictwa Narodowego, jak i innych, realizują liczne projekty animacji kulturowej na rzecz lokalnej publiczności. Projekty te można zakwalifikować jako projekty krótkie i o stosunkowo małym budżecie. Dzieje się to często w niepewnym i zmiennym otoczeniu prawnym i społecznym, przy jednoczesnym braku specjalistów w zakresie zarządzania projektami wewnątrz instytucji.

---

<sup>1</sup> Politechnika Wrocławska, Wydział Zarządzania, Katedra Systemów Zarządzania i Rozwoju Organizacji, e-mail: dorota.kuchta@pwr.edu.pl.

<sup>2</sup> Politechnika Wrocławska, Wydział Zarządzania, Katedra Systemów Zarządzania i Rozwoju Organizacji, e-mail: alicja.krawczynska@pwr.edu.pl.

Scrum jako sposób realizacji projektów odpowiada zarówno specyfice pracy małej instytucji kultury, jak i stałej konieczności adaptacji działań do zmieniających się warunków otoczenia. W oryginalnej postaci podejścia Scrum zarządzanie ryzykiem nie jest uwzględnione w sposób formalny, jednak w swojej istocie Scrum zawiera cykl zarządzania ryzykiem projektów.

W niniejszym rozdziale omówiono charakterystykę działalności instytucji kultury na szczeblu samorządowym w kontekście specyfiki podejmowanych projektów, następnie przedstawiono podstawowe pojęcia związane z zarządzaniem ryzykiem. Zaprezentowano także pojęcie Scruma jako pewnego rodzaju zbioru zasad dotyczącego prowadzenia projektów oraz zakres procesów zarządzania ryzykiem w ramach Scruma. Ostatnia część to opis autorskiego modelu zarządzania ryzykiem dostosowanego do potrzeb instytucji kultury z wykorzystaniem Scruma.

## **1.1. Projekty w małych instytucjach kultury w Polsce**

### **1.1.1. Instytucje kultury w Polsce na szczeblu samorządowym**

W Polsce funkcjonują 3 szczeble samorządu terytorialnego obejmujące 16 województw, 314 powiatów (w tym 66 miast na prawach powiatu) oraz 2477 gmin [GUS, 2019]. Prawa miejskie posiada 954 miejscowości, w tym 89% z nich to miasta poniżej 40 tysięcy mieszkańców oraz 59% to miasta do 10 tysięcy mieszkańców. Oznacza to, że oprócz dużych skupisk miejskich, takich jak Warszawa, Wrocław, Poznań, Kraków, Łódź oraz aglomeracji śląskiej i trójmiejskiej, życie kulturalne równolegle toczy się w małych miastach i miasteczkach. Nawet jeżeli obserwujemy coraz większą mobilność osób dorosłych i wewnętrzne procesy migracyjne, to nadal w znacznym stopniu świadomość kulturalna Polaków jest kształtowana poprzez małe ośrodki kulturalne na poziomie powiatu czy gminy.

Zasady organizacji i działalności instytucji kultury określa ustawa o organizowaniu i prowadzeniu działalności kulturalnej [1991]. Zgodnie z zapisami tej ustawy instytucjami kultury określa się instytucje państwowe oraz samorządowe, których podstawowym celem statutowym jest działalność kulturalna rozumiana jako tworzenie, upowszechnianie i ochrona kultury.

Według GUS w Polsce funkcjonuje ponad 13 tysięcy publicznych instytucji kultury, w tym 7881 bibliotek publicznych, 4255 ośrodków kultury, 959 muzeów oraz 188 teatrów i instytucji muzycznych [GUS, 2019]. Organ powołujący

instytucję kultury jest nazywany organizatorem i ma obowiązek zapewnienia środków niezbędnych do rozpoczęcia i prowadzenia działalności kulturalnej. Ponadto ustawa określa prowadzenie działalności kulturalnej jako zadanie własne jednostek samorządu terytorialnego o charakterze obowiązkowym.

Przychody instytucji kultury stanowią przede wszystkim dotacje podmiotowe i celowe. Dotacje podmiotowe zapewnia organizator na podstawie corocznych planów finansowych w danym roku budżetowym. Dotacje celowe mogą być przyznawane na szczeblu lokalnym lub krajowym w ramach programów Ministra Kultury, Dziedzictwa Narodowego i Sportu, jak również z innych publicznych i prywatnych źródeł (por. programy Narodowego Centrum Kultury, program Niepodległa, program Fundacji Banku Gospodarstwa Krajowego, Fundacji Orlen i inne). Dotacje celowe są przyznawane na realizację projektów w ramach procedur konkursowych. Przychody instytucji kultury mogą również pochodzić z prowadzonej działalności gospodarczej [Kowalik i in., 2011; Ilczuk, 2020].

### **1.1.2. Projekty realizowane przez instytucje kultury**

Zgodnie z definicją PMI (Project Management Institute) projekt to „tymczasowa działalność podejmowana w celu wytworzenia unikatowego wyrobu, dostarczenia unikatowej usługi bądź unikatowego rezultatu” [PMBOK, 2012, s. 4]. Poza celem i spodziewanym efektem, wskazanym w powyższej definicji, każdy projekt charakteryzuje określony termin (czas realizacji i zaplanowanie prac w czasie), budżet (uruchomione zasoby). Projekty mają charakter jednorazowy i tymczasowy [Larson, 2004; Pluszyńska, Konior, Gawęł, 2020]. Zatem projekt to przedsięwzięcie, które charakteryzuje celowość – projekt jest podejmowany, aby osiągnąć określony rezultat, tymczasowość – każdy projekt ma początek i koniec oraz wymóg posiadania zasobów – ludzi, środków finansowych, maszyn, obiektów itp. Zarządzanie projektem definiuje się jako „zastosowanie wiedzy, umiejętności, narzędzi i technik do działań projektowych w celu spełnienia wymagań projektu” [PMBOK, 2012, s. 5].

Projekty wpływają na procesy organizacyjne i artystyczne w zakresie kultury i sztuki [Ćwikła, 2016; Lindkvist, Hjort, 2015]. Zarządzanie projektem w instytucji kultury stało się elementem zarówno praktyki, jak i badań naukowych [Styhre, Börjesson, 2011; Malciene, Skaurone, 2017; Holder, 2018; Pluszyńska, Konior, Gawęł, 2020; Bemme, 2020]. W ramach dociekań naukowych autorzy analizują procesy zarządzania projektami w odniesieniu do instytucji

kultury, wskazując na analogie lub różnice z projektami realizowanym przez firmy komercyjne, a także na kierunki rozwoju zarządzania projektami.

W instytucji kultury projektem może być każde oryginalne i jednorazowe wydarzenie: koncert, wernisaż wystawy, spotkanie autorskie czy ciąg wydarzeń o jednej spójnej tematyce, związanej np. z promocją książki czy rocznicą urodzin Lema w Roku Lema. Mogą to również być warsztaty, plenery twórcze, biennale, pikniki i wiele innych działań wspierających na różne sposoby partycypację i animację kulturową społeczeństwa.

Specyfiką projektów w instytucji kultury w Polsce są małe projekty [Gołuchowski, Spyra, 2014; Ślusarczyk, 2016; Dziurski, Pawlicka, Wróblewska, 2017], zaś sposób finansowania ograniczają ich budżety o średniej około 70 tys. złotych i czas trwania do kilku miesięcy (średnia projektów w ramach programów MKDiNS z 2020 roku). Sposób realizacji projektów kulturalnych charakteryzują działania adaptacyjne bez wykorzystania metodyki, powtarzalność działań jako poczucie bezpieczeństwa, prowadzenie równocześnie wielu projektów w małych zespołach czy nawet przez jedną osobę [Pluszyńska, Konior, Gawel, 2020, s. 37]. Ponadto badania i praktyka wskazują z jednej strony na potrzebę profesjonalizacji tego obszaru funkcjonowania instytucji kultury [Dziurski, Pawlicka, Wróblewska, 2017], z drugiej zaś na projektyzację [Ćwikła, 2016], która oznacza nadmiar projektów we współczesnym świecie kultury.

## 1.2. Zarządzanie ryzykiem w projekcie

Projekt zawsze łączy się z przyszłością. Jest przedsięwzięciem jednorazowym i unikatowym, dlatego też zawsze towarzyszy mu ryzyko. Ryzyko w projekcie zostało zdefiniowane przez Pritcharda jako „skumulowany efekt prawdopodobieństwa niepewnych zdarzeń, które mogą korzystnie lub niekorzystnie wpływać na realizację projektu” [Pritchard, 2015, s. 7]. PMBOK ryzyko projektu określa nieco szerzej – jako „niepewne zdarzenie lub warunek, że jeśli wystąpi, ma pozytywny lub negatywny wpływ na jeden lub więcej celów projektu, takich jak zakres, harmonogram, koszt lub jakość” [PMBOK, 2012, s. 310]. Obie przytoczone definicje wskazują zarówno na negatywny, jak i pozytywny wpływ ryzyka na projekt. W literaturze przedmiotu i badaniach częściej wskazywany jest negatywny wpływ ryzyka na realizację celu projektu, a także na zakres, czas, budżet oraz jakość [Kuchta, Walczak, 2013; Andrzejewski, 2014; Juchniewicz, Metelski, 2017].

Zarządzanie ryzykiem w instytucjach kultury jest elementem procesu kontroli zarządczej. Zakres i sposób realizacji zarządzania ryzykiem został uregulowany Komunikatem Ministra Finansów „Szczegółowe wytyczne dla jednostek sektora finansów publicznych w zakresie planowania oraz zarządzania ryzykiem” [Ustawa o finansach publicznych, 2009; Komunikat Nr 6 Ministra Finansów, 2012]. Komunikat ten określa ryzyko jako możliwość zaistnienia zdarzenia, które negatywnie wpłynie na osiągnięcie celów i zadań. W związku z tym w dalszej części rozdziału ryzyko będzie rozpatrywane w kontekście negatywnego wpływu na realizację projektów.

O ile podejmowanie projektów jest zjawiskiem powszechnym, może mieć charakter intuicyjny, o tyle zarządzanie ryzykiem wiąże się z wiedzą zarówno domenową, jak i wiedzą dotyczącą zarządzania. W praktyce nie występują projekty, których realizacja nie jest warunkowana żadnymi wpływami zewnętrznymi i wewnętrznymi. Dlatego też jednym z podstawowych elementów zarządzania projektami, mającym wpływ na powodzenie całego przedsięwzięcia, jest proces zarządzania ryzykiem polegający na ograniczaniu i eliminowaniu niepożądanych zjawisk w projekcie [Andrzejewski, 2014]. Zarządzanie ryzykiem w projekcie jest jednym z ważniejszych aktywności w projekcie, ponieważ pozwala na osiąganie sukcesów niezależnie od występowania przypadkowych utrudnień [Voetsch, Cloffi, Anbari, 2004; Olechowski i in., 2016].

W literaturze przedmiotu można spotkać różne propozycje opisu procesu zarządzania ryzykiem zarówno na poziomie organizacji, jak i w projektach [Krysiak, Głowania, 2017; Muriana, Vizzini, 2017]. Pritchard zwraca uwagę, że zarządzanie ryzykiem w projekcie powinno się odbywać w perspektywie długoterminowej, w odniesieniu do całości projektu i do całej organizacji, jak również krótkoterminowej, w odniesieniu do najbliższego etapu w projekcie [Pritchard, 2015].

Na poziomie realizacji projektów zarządzanie ryzykiem jest najczęściej opisywane w formie procesu jako szereg uporządkowanych działań, na które składa się: identyfikacja czynników ryzyka, następnie oszacowanie skutków i prawdopodobieństwa zaistnienia ryzyka, opracowanie i wdrożenie planu reakcji na ryzyko wraz z budżetem oraz procesy kontroli zarządzania ryzykiem [Kuchta, Walczak, 2013; Hartono, 2018; Soroka-Potrzebna, 2020].

Według PMBOK [2012] zarządzanie ryzykiem projektu polega na identyfikowaniu, analizowaniu i reagowaniu na czynniki ryzyka przez cały okres trwania projektu w najlepszym interesie jego celów, zaś proces zarządzania ryzykiem w projekcie składa się z 6 następujących po sobie elementów:

1. Planowanie zarządzania ryzykiem.
2. Identyfikacja ryzyka.
3. Jakościowa analiza ryzyka.
4. Ilościowa analiza ryzyka.
5. Planowanie reakcji na ryzyko.
6. Monitorowanie i kontrolowanie ryzyka.

Metodyka zarządzania projektami PRINCE2™ [2009] określa elementy zarządzania ryzykiem w 4 krokach: identyfikacja, ocena, planowanie i wdrażanie, uzupełniając całość procesem komunikacji. Wskazuje również, że w trakcie trwania projektu pojawiają się nowe ryzyka, które mają wpływ zarówno na całość projektu, jak i na wcześniej zidentyfikowane ryzyka. Niezależnie od wskazanych powyżej elementów procesu zarządzania ryzykiem powinno ono prowadzić do zapewnienia założonych celów projektu.

### **1.3. Realizacja projektów z wykorzystaniem Scruma**

Scrum powstał z inicjatywy Kena Schwabera i Jeffa Sutherlanda, którzy w swoim opracowaniu pt. *Scrum Guide* definiują go jako „uproszczone ramy postępowania, które pomagają poszczególnym osobom, zespołom i organizacjom wytwarzać wartość poprzez adaptacyjne rozwiązywanie złożonych problemów” [www 3]. Podkreślają oni, że Scrum to pewien zestaw ogólnych reguł i narzędzi służący iteracyjnemu powstawaniu produktów. Zatem o wytwarzaniu wartości można mówić zarówno w kontekście produktu, jak i projektu. Chociaż Scrum powstał jako odpowiedź na potrzeby projektów IT, badania [Hall Walkey, 2015; Stoddard, Gillis, Cohn, 2019; Bemme, 2020] oraz raporty zastosowań Scruma [www 1] wskazują rosnące pozycje innych branż, w tym również instytucji publicznych.

Kluczowym elementem Scruma jest niewielki zespół (Scrum Team). To grupa osób mająca kompetencje do zrealizowania celu projektu. Scrum charakteryzuje się ponadto współpracą w ramach zespołu, klientocentrycznym podejściem i przejrzystością organizacyjną [Denning, 2015; Cockburn, 2016, Shutherland, 2017].

W Scrum cel projektu jest realizowany w sposób iteracyjny w ramach kolejnych sprintów. Są to wydarzenia o ustalonej długości trwające maksymalnie miesiąc. W trakcie tego okresu zespół dostarcza wartościowy przyrost produktu. Wartościowy przyrost oznacza przybliżanie realizacji celu projektu i pozwala na podejmowanie decyzji w odniesieniu do kolejnej iteracji. W ramach jednego sprintu występują 4 wydarzenia zwane ceremoniami. Są to: planowanie (pracy

na najbliższy sprint), codzienne spotkania zespołu, przegląd sprintu (podsumowanie pracy dotyczącej realizacji pracy w ramach sprintu) oraz retrospekcja (spotkanie zespołu, na którym analizuje się sposób pracy). Wydarzenia są powtarzalne w każdym kolejnym sprincie i wprowadzają uporządkowanie oraz regularność działań zespołu.

Scrum zakłada brak formalnej dokumentacji, którą zastępuje się artefaktami. Należą do nich:

- product backlog, czyli uporządkowana lista działań w projekcie; na górze listy znajdują się działania najważniejsze (najbardziej sensowne) do zrealizowania w najbliższym sprincie;
- sprint backlog, czyli lista zadań do zrealizowania w ramach jednego sprintu;
- przyrost, czyli cel i efekt każdego sprintu.

Scrum wyznacza także określone role w zespole, takie jak: Product Owner – odpowiedzialny przede wszystkim za rozwój produktu i zarządzanie backlogiem, Scrum Master – odpowiedzialny przede wszystkim za to, aby wszystkie procesy były realizowane zgodnie z regułami Scruma oraz Zespół Developerski, który jest odpowiedzialny (jako całość) za realizację celów sprintu. Zespół składa się maksymalnie z 9 osób. W założeniu zespół powinien być na tyle duży, aby gwarantować osiąganie celów w ramach sprintu i na tyle mały, aby pozostawać elastycznym oraz łatwo się ze sobą komunikować.

Scrum jako sposób realizacji projektów w swoich założeniach odpowiada specyfice pracy w małych instytucjach kultury. Przede wszystkim pozwala na rozpoczęcie projektu bez wszystkich danych wejściowych, jakim jest w pełni zdefiniowany zakres projektu kulturalnego, w przeciwieństwie do metod tradycyjnych kaskadowych [PRINCE2™, 2009; PMBOK, 2012]. Ponadto wskazuje na wagę małych zespołów [www 3], a takie właśnie zespoły realizują projekty w małych instytucjach kultury. Scrum pozwala na szybkie działanie i szybką adaptację do zmienności otoczenia [Sutherland, 2021], a taka była i jest nadal potrzeba w instytucjach kultury.

Istnieją nieliczne opracowania naukowe na temat wykorzystania Scruma w instytucjach kultury na świecie [Hall Walkey, 2015; Malciene, Skaurone, 2017; Stoddard, Gillis, Cohn, 2019], pozycje książkowe dotyczące zarządzania w kulturze z wykorzystaniem Scruma [Bemme, 2020] oraz relacje internetowe [www 2], które opisują jego zastosowania w kulturze.

## 1.4. Zarządzanie ryzykiem w Scrumie

Scrum został osadzony w teorii empiryzmu i reprezentuje pogląd, że wiedza wynika z doświadczenia, a podejmowanie decyzji bazuje na tym, co poznane. Przewodnik po Scrumie nie definiuje pojęcia ryzyka, natomiast bezpośrednio wskazuje, że „Scrum wykorzystuje iteracyjne, przyrostowe podejście w celu zwiększania przewidywalności oraz kontrolowania ryzyka” [www 3].

W Scrum Guide zarządzanie ryzykiem zostało opisane w dwóch wymiarach: przejrzystości i krótkich iteracji. Przejrzystość oznacza, że kształtujący się proces oraz praca muszą być widoczne dla osób ją wykonujących, jak również dla osób, na rzecz których praca ta jest wykonywana. W Scrumie decyzje są podejmowane na podstawie stanu jego trzech artefaktów (backlogu produktu, backlogu sprintu i przyrostu). Niedostateczna przejrzystość artefaktów może prowadzić do decyzji, których wynikiem jest zmniejszona wartość bądź zwiększone ryzyko. Czas trwania sprintu musi być odpowiednio dobrany. Kiedy sprint trwa zbyt długo, jego cel może się zdezaktualizować, może się zwiększyć złożoność, a ryzyko może wzrosnąć. Krótsze sprinty można wprowadzić dla uzyskania większej liczby cykli uczenia się oraz ograniczenia ryzyka związanego z kosztami i nakładem pracy do krótszych okresów. Każdy sprint można uznać za krótki projekt. Ponadto jednym z elementów działań w sprincie jest usuwanie przeszkód, czyli zdarzeń, które mają negatywny wpływ na realizację celu danego sprintu.

Dzięki tym podstawowym założeniom realizacja projektów w Scrumie przyczynia się do osiągnięcia założonych celów projektów [www 3], a przede wszystkim stałej adaptacji do zmieniających się warunków otoczenia i reakcji na ryzyko [Denning, 2015; Maximini, 2015; Moran, 2016].

## 1.5. Model zarządzania ryzykiem projektów z wykorzystaniem Scruma

Scrum nie wskazuje formalnych ram do zarządzania ryzykiem. Z przeglądu literatury i analizy prowadzonych badań w tym zakresie wynikają różne propozycje włączania aspektów procesu zarządzania ryzykiem w projektach realizowanych w Scrumie:

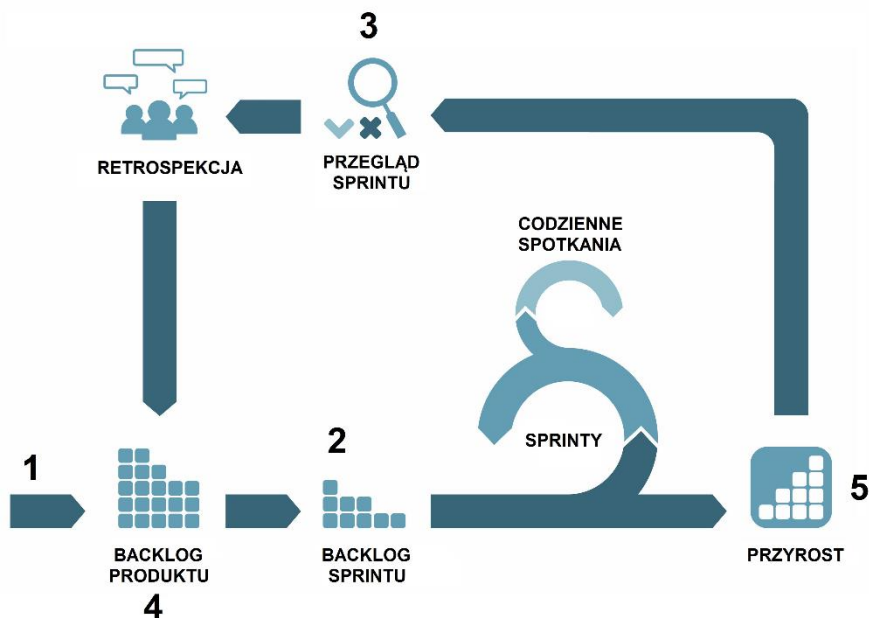
- połączenie Scruma i elementów kaskadowych metod zarządzania projektami, takich jak PMBOK [Chanouch, Mejri, Ghannouchi, 2019] i PRINCE2™ [Tomanek, Juricek, 2015];

- włączenie rejestru ryzyka do Scrum [Edzreena, Geer, Steward 2014; Chaouach, Mejri, Ghannouchi, 2019; Mousaei, Gandomani, 2018];
- analiza ryzyka podczas spotkań zespołu Scrum, Product Owner jako Risk Manager [Hammad, Inayat, 2018];
- uwzględnienie w backlogu produktu ryzyka pojedynczego zadania [Moran, 2016];
- połączenie kompotentów Scruma (wydarzeń, artefaktów, ról) i procesu zarządzania ryzykiem [Tavares i in., 2017, 2018, 2020].

Ponadto Alan Moran, profesor akademicki, a jednocześnie praktyk i konsultant w swojej książce *Agile Risk Management* [Moran, 2014] wskazuje, że zarządzanie ryzykiem powinno się odbywać w taki sam sposób, jak iteracjami i innymi elementami w ramach sprintów.

Mariusz Chrapko, praktyk w realizacji projektów, poleca prowadzenie rejestru przeszkód w celu zapobiegania im w przyszłych sprintach [Chrapko, 2015]. Ten sam autor proponuje, aby na etapie planowania prac projektowych opracować matrycę funkcjonalności projektu. W ramach tej matrycy określone są elementy ważne i nieważne dla celu projektu oraz takie, które mają wysokie lub niskie ryzyko realizacji, przy czym to ryzyko jest równoznaczne z trudnością i skomplikowaniem danego zadania.

Podsumowując: przedstawione pozycje literaturowe dotyczące zarządzania ryzykiem w Scrumie wskazują, że istnieje potrzeba włączania procesów zarządzania ryzykiem do procesów realizacji projektów realizowanych z wykorzystaniem Scruma. Procesy te mogą się uzupełniać. Ponadto dla większości z modeli prace badawcze realizowano w firmach dostarczających oprogramowanie, a statystyki jasno wskazują na rosnące zastosowanie Scruma w organizacjach i obszarach niezwiązanych z IT. Stąd pojawia się problem badawczy, jak dopasować istniejące metody zarządzania ryzykiem do projektów realizowanych w małych instytucjach kultury zarządzanych z wykorzystaniem Scruma. Przedstawiona poniżej autorska propozycja modelu zarządzania ryzykiem w projektach opartego na podejściu Scrum uwzględnia zarówno cechy projektów kulturalnych, jak i zarządzanie ryzykiem wynikające z procedur kontroli zarządczej w instytucjach kultury.



**Rys. 1.** Model zarządzania ryzykiem projektów z wykorzystaniem Scruma

Źródło: Opracowanie własne na podstawie [www 3].

W ramach modelu uwzględniono 5 komponentów łączących w sobie proces zarządzania ryzykiem i cykle Scruma. Proponowane komponenty zarządzania ryzykiem z wykorzystaniem Scruma wynikają z adaptacji procesu zarządzania ryzykiem i Scruma do specyfiki projektów kulturalnych.

### Komponent 1. Opracowanie backlogu projektu

Opracowanie backlogu poprzedza pierwszy sprint w projekcie. Scrum Guide używa określenia „backlog produktu”. Zastosowane określenie „backlog projektu” wynika z tego, że zastosowanie całego Scruma dotyczy projektu, a nie produktu. Dane wejściowe do backlogu stanowią następujące informacje:

- idea i zakres projektu (w przypadku projektów kulturalnych są to poszczególne wydarzenia i ich składowe);
- harmonogram wydarzeń w projekcie;
- działania, jakie należy podjąć, aby zrealizować poszczególne wydarzenia projektu, czyli pełna logistyka związana z realizacją projektu;
- doświadczenia z poprzednich projektów (np. sprawdzone działania promocyjne);
- wskaźniki rezultatu projektu.

Backlog jest listą uporządkowaną, co oznacza, że na samej górze znajdują się działania, które mają najwyższy priorytet i będą realizowane jako pierwsze. Główne kryteria, jakie należy wziąć pod uwagę przy określaniu priorytetów, obejmują: zaplanowane daty realizacji wydarzeń (np. datę koncertu czy wernisażu), dostępne zasoby osobowe i rzeczowe, złożoność i współzależność poszczególnych działań.

Ponadto określając zakres backlogu i priorytety, należy uwzględnić dane wynikające z dokumentów dotyczących kontroli zarządczej. Są one obowiązkowe w każdej instytucji publicznej i wynikają z Ustawy o finansach publicznych [2009]. Dokumentacja kontroli zarządczej obejmuje rejestr ryzyk, procedury zarządzania ryzykiem na poziomie organizacji, w tym identyfikację ryzyka, analizę ryzyka, określenie sposobu reakcji na ryzyko i monitoring. Dane te stanowią niezbędne elementy procesu zarządzania ryzykiem również na poziomie projektu.

Zarządzanie backlogiem należy do Product Ownera (specjalna rola w Scrumie), ale backlog jest otwarty i dostępny dla wszystkich członków zespołu projektowego. Mogą oni zatem dodawać działania i uzupełniać informacje w backlogu przed pierwszym sprintem oraz w czasie trwania całego projektu.

## **Komponent 2. Planowanie sprintu – tworzenie backlogu sprintu**

Każdy sprint rozpoczyna planowanie, którego celem jest wybór działań do zrealizowania w najbliższej iteracji czasowej. Zadania wybrane do sprintu trafiają do backlogu sprintu. W planowaniu uczestniczy cały zespół projektowy. Podczas planowania konieczne jest precyzyjne określenie, na czym będzie polegał przyrost oraz jakie efekty pracy będą oczekiwane na koniec sprintu. Omawiając zadania, należy wziąć pod uwagę przede wszystkim priorytety, czasochłonność zadań i dostępność zespołu realizacyjnego. Podczas tego spotkania konieczne jest przeanalizowanie czynników ryzyka poszczególnych zadań, sposobu reakcji w przypadku wystąpienia takiego ryzyka. Podstawowym dokumentem stanowiącym wytyczne w tym zakresie jest rejestr ryzyk w organizacji.

## **Komponent 3. Przegląd sprintu**

W przeglądzie sprintu uczestniczą wszyscy członkowie zespołu projektowego (Scrum Team) zarządzający instytucją, interesariusze oraz inne osoby, które są powiązane z projektem. Spotkanie rozpoczyna się prezentacją efektów pracy zespołu z danego sprintu. W czasie prezentacji członkowie zespołu projektowego omawiają również przeszkody i sytuacje negatywne, które miały wpływ na projekt. Jeżeli te przeszkody mają wpływ na dalszy przebieg projektu, należy

je uwzględnić w backlogu projektu. Ostatnim krokiem jest formalna akceptacja przyrostu lub wskazanie zmian, które należy uwzględnić w kolejnym sprincie.

#### **Komponent 4. Aktualizacja i zarządzanie backlogiem**

Aktualizacja backlogu obejmuje stałe uzupełnianie listy o nowe informacje na podstawie już zrealizowanych części projektu oraz informacje płynące z otoczenia zewnętrznego projektu i całej instytucji. Aktualizacja backlogu polega także na redefiniowaniu priorytetów działań, co jest szczególnie ważne wobec zmieniającego się otoczenia instytucji. W ramach zarządzania backlogiem konieczna jest stała identyfikacja czynników ryzyka i planowanie reakcji na ryzyko.

#### **Komponent 5. Iteracyjny przyrost**

Przyrost produktu to zrealizowane w sprincie działania, które wchodzą w zakres całego projektu. Przyrostem jest wszystko, co przynosi wartość. W instytucji kultury będą to np.: pozyskanie partnerów medialnych, ustalenie warunków kontraktu z artystą, opracowanie planu promocji wydarzenia czy zatwierdzenie koncepcji graficznej wydarzenia. Przyrost występuje zatem zarówno na poziomie jednego sprintu, jak i na poziomie projektu. Dlatego też, poza widocznymi efektami projektu, jest to wiedza członków zespołu projektowego i ich doświadczenie. Do przyrostu zalicza się także opracowana dokumentacja (np. umowa z wykonawcą, która może być również wykorzystana jako wzór w kolejnych projektach) czy nowe procedury działania (np. sposób zlecania zadań podwykonawcom na zewnątrz instytucji). W odniesieniu do zarządzania ryzykiem przyrost powinien obejmować także aktualizację rejestru ryzyk na poziomie organizacji.

### **Podsumowanie**

Jak już zostało wspomniane w pierwszej części niniejszego rozdziału, projekty realizowane przez instytucje kultury charakteryzuje krótki czas trwania (od 3-9 miesięcy) oraz średnie budżety na poziomie około 60-70 tysięcy złotych. Ponadto instytucje podejmują wiele projektów w tym samym czasie. Realizacją projektów zajmują się małe zespoły, które działają adaptacyjnie, bez wykorzystania metodyki, a wiedzę czerpią z doświadczenia. Powtarzalność działań wzmacnia poczucie bezpieczeństwa realizatorów.

Model zarządzania ryzykiem projektów w małych instytucjach kultury z wykorzystaniem Scruma bazuje na cykliczności procesu zarządzania ryzykiem w projektach, zaczynając od identyfikacji, przez ocenę, reakcję na ryzyko, a na

monitorowaniu ryzyka kończąc. Równocześnie został wpisany w cykliczność i iteracyjność sprintów w Scrumie. Zaproponowany model zostanie poddany analizie i walidacji podczas badań pilotażowych realizowanych w formie wielokrotnego studium przypadku w małych instytucjach kultury.

## Literatura

- A Guide to the Project Management Body of Knowledge (PMBOK)* (2012), 5th Edition, Project Management Institute.
- Andrzejewski M. (2014), *Ryzyko w zarządzaniu projektami* [w:] M.J. Szymankiewicz, P. Kuźbik (red.), *Zarządzanie organizacją z perspektywy metodologicznej*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Bemme S.O. (2020), *Kultur-Projektmanagement. Kultur- und Organisationsprojekte erfolgreich managen*, Springer VS, Wiesbaden.
- Chaouch S., Mejri A., Ghannouchi S.A. (2019), *A Framework for Risk Management in Scrum Development Process*, "Procedia Computer Science", Vol. 164, s. 188-192.
- Chrapko M. (2015), *Scrum. O zwinnym zarządzaniu projektami*, Helion, Gliwice.
- Cockburn A. (2016), *The Heart of Agile. Humans and Technology Technical Report 2016*, Florida.
- Ćwikła M. (2016), *Projekt to jest projekt. Specyfika zarządzania projektami kulturalnymi na przykładzie tworzenia koprodukcji teatralnych*, Attyka, Kraków.
- Ćwikła M. (2017), *Cultural Projects in 2030: A Performative Approach* [w:] *ReThinking Management*, Springer VS, Wiesbaden.
- Denning S. (2015), *Agile: It's Time to Put It to Use to Manage Business Complexity*, Strategy & Leadership, Bingley.
- Dziurski P., Pawlicka K., Wróblewska A. (2017), *Wokół zarządzania kulturą. Rezultaty badania ankietowego*, Załącznik Kulturoznawczy, Uniwersytet Kardynała Stefana Wyszyńskiego, Warszawa.
- Edzreena O., Des G., Stewart D. (2014), *Lightweight Risk Management in Agile Projects*, 26th Software Engineering Knowledge Engineering Conference (SEKE), Vol. 26, s. 576-581.
- Gołuchowski J., Spyra Z. (2014), *Zarządzanie w kulturze, sztuce i turystyce kulturowej*, CeDeWu, Warszawa.
- GUS (2019), *Rocznik Statystyczny Głównego Urzędu Statystycznego*, Warszawa.
- Hall Walkey L. (2015), *Changing the Workplace Culture at Flinders University Library* [w:] *From Pragmatism to Professional Reflection*, Australian Academic & Research Libraries, Sydney.

- Hammad M., Inayat I. (2018), *Integrating Risk Management in Scrum Framework*, "International Conference on Frontiers of Information Technology (FIT)", s. 158-163.
- Hartono B. (2018), *From Project Risk to Complexity Analysis: A Systematic Classification*, "International Journal of Managing Projects in Business", Vol. 11, Bingley, s. 734-760.
- Holder S. (2018), *Borrowed from Business: Using Corporate Strategies to Manage Library Projects* [w:] *Project Management in the Library Workplace*, Bingley.
- Ilczuk D. (2020), *System finansowania instytucji kultury w Polsce. Opracowanie Senatu RP*, Centrum Badań nad Gospodarką Kreatywną SWPS Uniwersytet Humanistyczno-Społeczny, Warszawa.
- Juchniewicz M., Metelski W. (2017), *Zarządzanie ryzykiem w projektach – kontekst behawioralny* [w:] J. Szumniak-Samolej (red.), *Współczesne wyzwania w zakresie funkcjonowania przedsiębiorstw. Perspektywa badawcza młodych naukowców*, Oficyna Wydawnicza SGH, Warszawa.
- Komunikat Nr 23 Ministra Finansów z dnia 16 grudnia 2009 r. w sprawie ogłoszenia „Standardów kontroli zarządczej w jednostkach sektora finansów publicznych” (Dz.Urz. M.F. Nr 15, poz. 84).
- Komunikat Nr 6 Ministra Finansów z dnia 6 grudnia 2012 r. w sprawie szczegółowych wytycznych dla sektora finansów publicznych w zakresie planowania i zarządzania ryzykiem (Dz.Urz. M.F. z 2012 r., poz. 56).
- Kowalik W., Matlak M., Nowak A., Noworól K., Noworól Z. (2011), *Kultura lokalnie. Między uczestnictwem w kulturze a partycypacją w zarządzaniu*, Małopolski Instytut Kultury, Kraków.
- Krysiak M., Głowania S. (2017), *Metodyki zarządzania projektami IT i ich ryzykiem: przegląd i wykorzystanie*, „Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 340, Katowice, s. 79-98.
- Kuchta D., Walczak W. (2013), *Risk Characteristic of Agile Project Management Methodologies and Responses to Them*, "Operations Research and Decisions", Vol. 4, s. 75-95.
- Larson R. (2004), *The Critical Steps to Managing Small Projects*, PMI® Global Congress, Praga.
- Lindkvist L., Hjort D. (2015), *Organizing Cultural Project Through Legitimising as Cultural Entrepreneurship*, "International Journal of Managing Project in Business", Bingley, s. 696-714.
- Malciene Z., Skaurone L. (2017), *Project Method Management Organizing Cultural Events*, „Taikomiejai tyrimai studijose ir praktikoje – Applied Research in Studies and Practice”, Vol. 13, Wilno, s. 48-54.
- Maximini D. (2015), *The Scrum Culture Introducing Agile Methods in Organizations*, Springer International Publishing, New York.
- Moran A. (2014), *Agile Risk Management*, Berlin, Springer.

- Moran A. (2016), *Risk Management in Agile Projects*, "Isaca Journal", Vol. 2, Schaumburg, s. 1-4.
- Mousaei M., Gandomani T.J. (2018), *A New Project Risk Management Model Based on Scrum Framework and Prince2 Methodology*, "International Journal of Advanced Computer Science and Applications", Vol. 4, s. 442-449.
- Muriana C., Vizzini G. (2017), *Project Risk Management: A Deterministic Quantitative Technique for Assessment and Mitigation*, "International Journal of Project Management", Vol. 35, Wiedeń, s. 320-340.
- Olechowski A., Oehmen J., Seering W., Ben-Daya M. (2016), *The Professionalization of Risk Management: What Role Can the ISO 31000 Risk Management Principles Play?* "International Journal for Project Management", Vol. 34, Wiedeń, s. 1568-1578.
- Pluszyńska A., Konior A., Gawel Ł. (2020), *Zarządzanie w kulturze, teoria i praktyka*, WN PWN, Warszawa.
- PRINCE2™ (*Projects in Controlled Environments*) (2009), The Stationery Office (TSO) Managing Successful Projects with PRINCE2™.
- Pritchard C.L. (2015), *Risk Management. Concepts and Guidance*, Fifth Edition, CRC Press, Taylor & Francis Group, Boca Raton.
- Skorupka D., Kuchta D., Górski M. (2012), *Zarządzanie ryzykiem w projekcie*, Wyższa Szkoła Oficerska Wojsk Lądowych im. generała Tadeusza Kościuszki, Wrocław.
- Soroka-Potrzebna H. (2020), *Identyfikacja ryzyka w zarządzaniu projektami – praktycy a eksperci* [w:] E. Sońta-Drączkowska, I. Bednarska-Wnuk (red.), *Wybrane aspekty zarządzania procesami i projektami w przedsiębiorstwach*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Stoddard M.M., Gillis B., Cohn P. (2019), *Agile Project Management in Libraries: Creating Collaborative, Resilient, Responsive Organizations*, The Journal of Library Administration, The George Washington University.
- Styhre A., Börjesson S. (2011), *Project Management in the Culture Industry: Balancing Structure and Creativity*, "International Journal of Project Organisation and Management (IJPOM)", Vol. 3, Geneva, s. 22-35.
- Sutherland J. (2021), *Scrum, czyli jak robić dwa razy więcej, dwa razy szybciej*, Witkom Salma Press, Warszawa.
- Ślusarczyk A. (2016), *Determinanty skutecznego zarządzania projektami kultury*, „Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 256, Katowice, s. 67-80.
- Tavares B.G., Sanches da Silva C., Diniz de Souza A. (2017), *Risk Management in Scrum Software Projects*, "International Transactions in Operational Research", Vol. 26, s. 1884-1905.
- Tavares B.G., Silva C.E.S, Souza A.D. (2018), *Practices to Improve Risk Management in Agile Projects*, "International Journal of Software Engineering and Knowledge Engineering", Vol. 2, Singapur, s. 1-19.

- Tavares B.G., Silvia C.E.S., Souza A.D., Keil W. (2020), *A Risk Management Tool for Agile Software Development*, "Journal of Computer Information Systems", Vol. 6, Philadelphia, s. 561-570.
- Tomanek M., Juricek J. (2015), *Project Risk Management Model Based on PRINCE2™ and Scrum Frameworks*, "International Journal of Software Engineering & Applications (IJSEA)", Vol. 1, Sydney, s. 81-88.
- Ustawa z dnia 25 października 1991 r. o organizowaniu i prowadzeniu działalności kulturalnej (Dz.U. 1991 Nr 114, poz. 493 z późn. zm.).
- Ustawa z dnia 27 sierpnia 2009 r. o finansach publicznych. Dz.U. Nr 157, poz. 1240 z późn. zm.
- Voetsch R.J., Cloffi D.F., Anbari F.T. (2004), *Project Risk Management Practises and Their Association with Reported Project Success*, © IRNOP VI Conference, Turku.
- Zarządzanie ryzykiem. Informacje ogólne (2011), Departament Audytu Sektora Finansów Publicznych, Ministerstwo Finansów, Warszawa.
- [www 1] 14<sup>th</sup> Annual State of Agile Raport, 2020, <https://www.stateofagile.com/> (dostęp: 20.04.2021).
- [www 2] <https://handschriftenportal.de/scrum-in-der-bibliothek/> (dostęp: 20.04.2021).
- [www 3] Scrum Guide 2020, [www.scrum.org](http://www.scrum.org) (dostęp: 20.04.2021).



## **Sprzężenie zwrotne w adaptacyjnym zarządzaniu ryzykiem projektu**

*Dariusz Meiser<sup>1</sup>*

### **Wprowadzenie**

W rozdziale przedstawiono koncepcje „multisprzężenia zwrotnego” oraz „samosprzężenia zwrotnego”, które stanowią część prac badawczych autora nad metodą adaptacyjnego zarządzania ryzykiem projektu wytwórczego. Celem tego rozdziału jest wstępne i ogólne nakreślenie idei proponowanych sprzężeń zwrotnych, których szczegółowe aspekty rozwoju, integracji z metodą adaptacyjnego zarządzania projektami oraz przykłady praktycznej aplikacji będą przedmiotem dalszych prac autora.

Rozdział składa się z trzech części oraz wprowadzenia i podsumowania. W pierwszej części omówiono pojęcie zarządzania ryzykiem oraz typowe fazy zarządzania ryzykiem według najczęściej stosowanych metod zarządzania projektami, tj. PMI (PMBOK) oraz PRINCE2. Kolejny fragment dotyczy adaptacyjnego procesu zarządzania ryzykiem projektu według propozycji autora, której uzupełnieniem jest autorska koncepcja sprzężeń zwrotnych zaprezentowana w części trzeciej rozdziału.

Realizacja każdego projektu jest nieodłącznie związana z występującym ryzykiem. Proces zarządzania ryzykiem projektu powinien zatem stanowić integralną część procesu zarządzania samym projektem. Te dwa procesy zarządcze nie mogą być realizowane obok siebie czy też niezależnie od siebie – w odosobnieniu. W procesie zarządzania ryzykiem projektu należy wyraźnie odróżnić zmniejszanie możliwości (prawdopodobieństwa) wystąpienia danego ryzyka od ograniczania jego skutków w sytuacji, gdy dane ryzyko rzeczywiście wystąpi. Reaktywne zarządzanie ryzykiem sprowadza się wyłącznie do tego drugiego aspektu – minimalizacji niepożądanych skutków wystąpienia danego ryzyka.

---

<sup>1</sup> Uniwersytet Ekonomiczny w Katowicach, Wydział Informatyki i Komunikacji, Katedra Badań Operacyjnych, e-mail: meiser.consulting@gmail.com.

Zarządzanie proaktywne skupia się natomiast na obydwu wątkach [Smith, 2002, s. 4]: ograniczaniu możliwości, że dane ryzyko wystąpi, oraz redukcji skutków jego wystąpienia w momencie, gdy dane ryzyko rzeczywiście się zmaterializuje.

## 2.1. Zarządzanie ryzykiem projektu – ujęcie standardowe

Jako standardowe ujęcie zarządzania ryzykiem projektu autor określa klasyczne metody zarządzania projektami zaliczane do modeli TPM (ang. *Traditional Project Management*) – tradycyjnego zarządzania projektami [Wysocki, 2013, s. 79, 80]. Ma ono zastosowanie w sytuacjach, gdy zarówno cel, jak i rozwiązanie projektu są znane i dokładnie zdefiniowane przed rozpoczęciem prac projektowych. Tego typu projekty były w przeszłości wielokrotnie realizowane, wobec czego wykonujące je zespoły mają bogate doświadczenie oraz wypracowane schematy postępowania podczas ich wykonywania. Infrastruktura, warunki realizacji oraz otoczenie projektu są również znane, sprawdzone i przewidywalne. W literaturze dotyczącej zarządzania projektami taki model jest nazywany często *Plan-Driven Method* [Boehm, Turner, 2012, s. 9] ze względu na fakt, że projekty TPM są projektami opartymi „na planowaniu i realizacji planów” [Wysocki, 2013, s. 88], a ich sukces jest postrzegany „przez pryzmat zgodności z opracowanym planem” [Wysocki, 2013, s. 89]. W projektach klasy TPM zakłada się, że nie będzie zmian związanych z zakresem, wymaganiami oraz przebiegiem realizacji takiego projektu.

Aby w ogóle rozpocząć rozważania nad różnymi technikami zarządzania ryzykiem w projektach, należałoby zdefiniować samo pojęcie zarządzania ryzykiem. Na początek przywołajmy definicję zawartą w normie PN-ISO 31000 „Zarządzanie ryzykiem. Zasady i wytyczne”. Według niej zarządzanie ryzykiem jest określone jako „skoordynowane działania dotyczące kierowania i nadzorowania organizacją w odniesieniu do ryzyka” [PN-ISO 31000:2012, s. 17].

Cytowana definicja zarządzania projektami z normy PN-ISO 31000 jest bardzo zwięzła, lecz w codziennej praktyce projektowej zbyt ogólna do zastosowania. Warto się zatem przyjrzeć, co na ten temat mówią najbardziej rozpowszechnione metody zarządzania projektami, tj. PMI (PMBOK) oraz PRINCE2. Te dwie metody zostały przez autora wybrane ze względu na ich dominującą rolę w zarządzaniu projektami. Wskazują na to między innymi badania przeprowadzone przez Jasińską i Szapiro na 75 respondentach z polskich firm realizujących projekty ICT. Badania te dotyczyły popularności stosowania różnych me-

toż zarządzania projektami w tych firmach i wykazały, że: 31% respondentów wskazało na metodę PMBOK, 16% na metodę PRINCE2, a 9% zadeklarowało stosowanie metody zwinnej SCRUM [Jasińska, Szapiro, 2014, s. 208] (pozostałe 44% badanych odpowiedziało, że „stosuje własne metodyki zarządzania projektem”). Z kolei badania (za okres 2005-2012) przytoczone przez Trockiego wykazały, że według statystyk APMG<sup>2</sup> w zakresie popularności metody PRINCE2 Polska zajmuje 5. miejsce na świecie „pod względem liczby wydanych certyfikatów znajomości tej metodyki na poziomie Foundation (23 058 osób) oraz Practitioner (3676 osób)” [Trocki, 2017, s. 145]. Niepublikowane badania autora wśród polskich, tajwańskich, niemieckich oraz szwedzkich firm realizujących procesy zarządzania projektami również wskazują na dominującą rolę metod PMI (PMBOK) oraz PRINCE2.

W metodzie PMI przedstawiono sześćoetapowy proces zarządzania ryzykiem [PMBOK Guide, 2000, s. 127-146]:

- **planowanie** zarządzania ryzykiem (ang. *Risk Management Planning*),
- **identyfikacja ryzyka** (ang. *Risk Identification*),
- **jakościowa analiza** ryzyka (ang. *Qualitative Risk Analysis*),
- **ilościowa analiza** ryzyka (ang. *Quantitative Risk Analysis*),
- **planowanie reakcji** na ryzyko (ang. *Risk Response Planning*),
- **monitorowanie i kontrola** ryzyka (ang. *Risk Monitoring and Control*).

Proces ten zobrazowano na rys. 1.



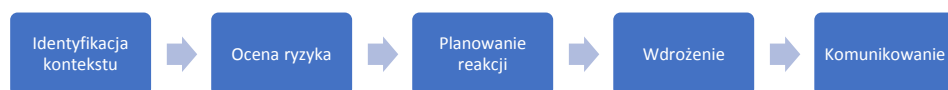
**Rys. 1.** Proces zarządzania ryzykiem projektu według PMI (PMBOK)

Z kolei do zarządzania ryzykiem w metodzie PRINCE2 służy narzędzie „Management of Risk” (MoR) oparte na pięcioetapowym procesie zobrazowanym na rys. 2 [Wyrozębski, 2011, s. 106-107]:

- **identyfikacja kontekstu** projektu oraz ryzyka projektowego i jego umieszczenie w rejestrze ryzyka,
- **ocena ryzyka**, tj. szacowanie prawdopodobieństwa i skutków oraz ewaluacja łącznej wartości ryzyka w projekcie,

<sup>2</sup> APMG (APM Group) International jest jednym z wiodących globalnych instytutów akredytacyjnych i egzaminacyjnych.

- **planowanie** szczegółowych **reakcji** na ryzyka w celu usunięcia bądź redukcji zagrożeń i maksymalizacji szans, w tym także uwzględnienie tych reakcji w planie projektu, planie etapu, uzasadnieniu biznesowym projektu,
- **wdrożenie** zaplanowanych działań, monitorowanie ich efektów oraz podejmowanie stosownych działań korygujących,
- **komunikowanie** ryzyka interesariuszom (wewnętrznym i zewnętrznym) projektu; komunikowanie jest zadaniem ciągłym i odbywa się zazwyczaj z wykorzystaniem raportów z punktów kontrolnych, raportów nadzwyczajnych, raportów końcowych etapu i raportów końcowych projektu.



**Rys. 2.** Proces zarządzania ryzykiem projektu według PRINCE2 (MoR – *Management of Risk*)

Analizując również inne występujące w literaturze definicje i opisy procesów zarządzania ryzykiem w projektach, można sformułować spostrzeżenie, że typowy proces zarządzania ryzykiem składa się z kilku podstawowych czynności, które są wymieniane praktycznie we wszystkich metodach i narzędziach. Zaliczają się do nich takie działania, jak: identyfikacja, analiza, zaplanowanie działań, wykonanie działań, monitorowanie skuteczności przeprowadzonych działań oraz przebiegu samego procesu. Można zatem przyjąć, że są to działania konieczne, aczkolwiek niewystarczające, do skutecznej realizacji procesu zarządzania ryzykiem w projektach. Szczegółowe przedstawienie takiej analizy przeprowadzonej przez autora w ramach prowadzonych prac badawczych wykracza poza ramy tego rozdziału.

## 2.2. Adaptacyjne zarządzanie ryzykiem projektu

Adaptacyjność w zarządzaniu ryzykiem projektu można rozumieć jako ciągłe przystosowywanie się – adaptację – do zmieniających się warunków, w jakich jest realizowany projekt w momencie, gdy wystąpią jakiegokolwiek zmiany mogące mieć wpływ na ryzyko projektu.

W przeciwieństwie do reaktywnego, pasywnego zarządzania podejście adaptacyjne kładzie nacisk na ciągłe monitorowanie otoczenia i wyszukiwanie ewentualnych zmian, a także bieżącą aktualizację ryzyk oraz opracowywanie

i modyfikację planów zapobiegawczych czy awaryjnych. Taki sposób działania symbolizuje zamknięta pętla schematu blokowego adaptacyjnego zarządzania ryzykiem w ujęciu zaproponowanym przez autora, przebiegająca w zobrazowany poniżej sposób.



**Rys. 3.** Adaptacyjny proces zarządzania ryzykiem projektu – propozycja autora

Poszczególne fazy adaptacyjnego procesu zarządzania ryzykiem projektu składają się z następujących działań:

1. **Faza eksploracji:** poszukiwanie ryzyk, identyfikacja ryzyk.
2. **Faza operacjonalizacji:** analiza ryzyk, szacowanie prawdopodobieństwa i wpływu ryzyk na projekt, opracowanie planów: zapobiegawczych i awaryjnych.
3. **Faza zarządzania:** bieżące zarządzanie ryzykami, aktualizacja ryzyk, ciągłe monitorowanie ryzyk, retrospekcja.

### Faza eksploracji

Działania w ramach pierwszej fazy – eksploracji, mają na celu poszukiwanie oraz identyfikację ryzyk występujących w danym projekcie. Za każdym razem przy rozpoczęciu nowego projektu wejście do rejestru zarządzania ryzykiem stanowi „skumulowany rejestr ryzyk” dla projektów wykonywanych

w przeszłości. Jako dodatkowy czynnik określający poziom ryzyka – oprócz „prawdopodobieństwa wystąpienia” oraz „siły wpływu na projekt” – autor proponuje przyjąć miarę, która będzie określała, czy dane ryzyko już kiedyś wystąpiło w jednym z poprzednich projektów. Taka miara określa procentowo, w ilu poprzednich projektach dane ryzyko wystąpiło i – według propozycji autora – jest określana jako: „**współczynnik częstości ryzyka**” (%). Na przykład dla ryzyka „problemy z EMC” współczynnik oblicza się jako: na 100 zrealizowanych projektów problemy z EMC wystąpiły 26 razy, co oznacza, że współczynnik częstości dla tego ryzyka to 26%. Następnie wszystkie ryzyka są sortowane względem współczynnika częstości, zaczynając od najwyższych, ponieważ to one wskazują na ryzyka, z którymi w firmie najczęściej występują problemy. Takie ryzyka zazwyczaj powtarzają się z projektu na projekt, co oznacza, że z tego typu ryzykami dana organizacja ma poważny problem i niczego się tak naprawdę w tym zakresie nie nauczyła. Dlatego też w każdym nowym projekcie należy zwrócić szczególną uwagę na te ryzyka, które w poprzednich projektach najczęściej przekształcały się z „ryzyk” w „incydenty”. W kolejnym kroku wśród ryzyk z najwyższym współczynnikiem częstości szczególną uwagę trzeba zwrócić na te ryzyka, które mają najwyższe parametry „siły wpływu” na projekt oraz „prawdopodobieństwa wystąpienia”.

Proces opracowywania współczynnika częstości ryzyka oraz skumulowanego rejestru ryzyk w postaci usystematyzowanej procedury będzie przedmiotem dalszych prac autora.

## **Faza operacjonalizacji**

Celem pierwszej fazy – eksploracji, było znalezienie oraz zidentyfikowanie ryzyk charakterystycznych dla danego projektu. Kolejna faza – operacjonalizacji, składa się z trzech następujących po sobie kroków: analizy ryzyk, szacowania prawdopodobieństw i wpływu ryzyk na projekt oraz opracowania planów: zapobiegawczych i awaryjnych. Głównym celem operacjonalizacji jest wykonanie analizy i podjęcie decyzji o tym, które z ryzyk zidentyfikowanych w fazie eksploracji są na tyle znaczące, aby się nimi aktywnie zajmować. W procesach zarządzania ryzykiem realizowanych według modelu liniowego (TPM) często zdarza się, że po przeprowadzeniu identyfikacji ryzyk jest opracowywana długa lista ryzyk i wszystkie one intencjonalnie mają podlegać aktywnemu zarządzaniu przez zespół projektowy. Jednak najczęściej bywa tak, że po rozpoczęciu projektu zespół koncentruje się na rzeczywistych pracach nad projektem, pomijając kwestie związane z zarządzaniem ryzykiem [Smith, Merritt, 2002, s. 61].

Nie jest bowiem możliwe skuteczne zarządzanie kilkudziesięcioma ryzykami zidentyfikowanymi dla danego projektu, gdzie każde z nich może dotyczyć innego obszaru, różnych zagadnień, a przede wszystkim może mieć skrajnie różny wpływ na osiągnięcie celów danego projektu. Dlatego też kluczową kwestią jest przeprowadzenie dokładnej analizy wcześniej zidentyfikowanych ryzyk i wyodrębnienie tych, które będą aktywnie zarządzane.

Dla każdego zidentyfikowanego w poprzedniej fazie ryzyka należy odpowiedzieć na następujące pytania [Smith, Merritt, 2002, s. 61]:

- dlaczego – zdaniem analizującego – dane ryzyko może wystąpić?
- jaka będzie całkowita strata, gdy to ryzyko wystąpi?
- skąd wniosek, że całkowita strata będzie miała konkretną wartość?
- jak są szacowane subiektywne prawdopodobieństwa: wystąpienia ryzyka oraz jego wpływu na projekt?

## **Faza zarządzania**

Na ostatnią fazę – zarządzania – składają się następujące działania: bieżące zarządzanie ryzykami, aktualizacja ryzyk, ciągłe monitorowanie ryzyk oraz retrospekcja.

### **Bieżące zarządzanie ryzykami**

Jest ono realizowane na podstawie opracowanych w fazie operacjonalizacji planów zapobiegawczych i planów awaryjnych. Plany zapobiegawcze są przygotowywane w celu powstrzymania zdarzenia wywołującego ryzyko; plany awaryjne oraz rezerwy – w celu zminimalizowania straty w przypadku, gdy dane ryzyko rzeczywiście wystąpi.

### **Ciągłe monitorowanie ryzyk**

Regularny przegląd i monitorowanie ryzyka ma trzy cele:

- czy działania podjęte w celu ograniczania możliwości wystąpienia ryzyka lub ograniczania jego skutków w przypadku, gdy dane ryzyko rzeczywiście wystąpi, przynoszą zakładane efekty?
- czy warunki – wewnętrzne lub zewnętrzne – mające wpływ na zidentyfikowane ryzyka uległy zmianie?
- czy wystąpiła możliwość pojawienia się nowych, wcześniej niezidentyfikowanych ryzyk?

## Retrospekcja

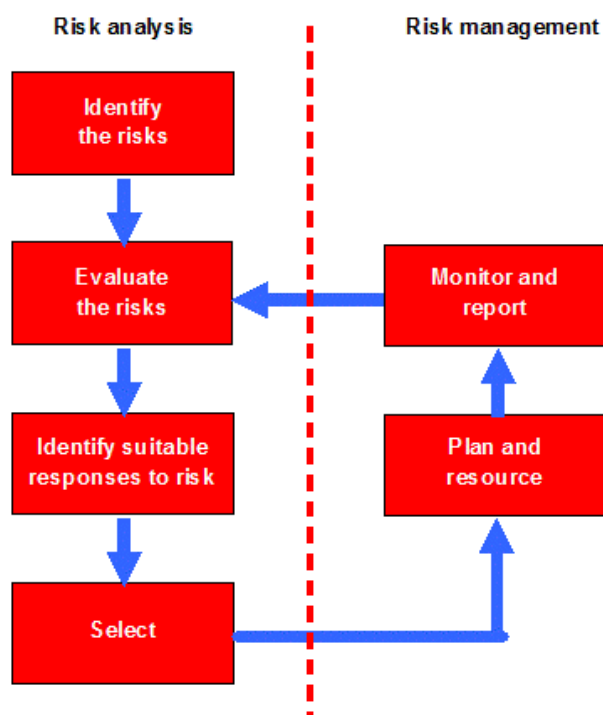
Retrospekcja stanowi ostatnie działanie w ramach fazy zarządzania. Po zakończeniu projektu – niezależnie od powodu, dla którego ono nastąpiło – trzeba koniecznie skonfrontować (historycznie) „prawdopodobieństwo wystąpienia” danego ryzyka, które było określone przed realizacją danego projektu ze „współczynnikiem częstości”, czyli jak często dane ryzyko rzeczywiście wystąpiło, przekształcając się w incydent. Tak samo jak przy zarządzaniu projektem jednym z ważniejszych zadań jest zamykanie projektu – spotkanie zamykające + analiza *post factum* – tak w zarządzaniu ryzykiem projektu jedną z ważniejszych kwestii jest przeprowadzenie retrospekcji.

Przedstawiony model adaptacyjnego zarządzania projektem jest z założenia uniwersalny i nie zależy od wyboru określonej metody zarządzania projektem ani od jego cech charakterystycznych. Może być on używany praktycznie dla każdego modelu cyklu życia projektu: TPM (ang. *Traditional Project Management*) – model tradycyjnego zarządzania projektami, APM (ang. *Agile Project Management*) – model zwinnego zarządzania projektem, xPM (ang. *Extreme Project Management*) – model ekstremalny „klasyczny”, MPx (ang. *Emertxe Project Management*) – model ekstremalny „emertxe” [Wysocki, 2013, s. 79, 80, 84]. Integracja zaproponowanej przez autora adaptacyjnej metody zarządzania ryzykiem z modelami TPM, APM, xPM oraz MPx będzie przedmiotem dalszych jego prac.

## 2.3. Koncepcja sprzężenia zwrotnego

Schemat procesu zaproponowany przez autora na rys. 3 uwzględnia już powtarzalny, iteracyjny cykl działania. Jednak brakuje w nim jednego istotnego elementu – sprzężenia zwrotnego. Aby przedstawić koncepcje proponowane przez autora w tym zakresie, trzeba zacząć od nakreślenia pojęcia „sprzężenia zwrotnego”. W najbardziej ogólnym ujęciu systemowym sprzężenie zwrotne polega na doprowadzeniu części sygnału wyjściowego z powrotem do wejścia. Część sygnału wyjściowego, zwana sygnałem zwrotnym, zostaje skierowana do wejścia układu i zsumowana z sygnałem wejściowym, wskutek czego zmieniają się warunki sterowania układu. Inaczej mówiąc, jest to mechanizm oddziaływania w formie bezpośredniej lub pośredniej zmian na wyjściach danego systemu na stan jego wejść. Tak więc ideą działania sprzężenia zwrotnego jest dostosowanie kolejnych odpowiedzi systemu na podstawie informacji dotyczących efektów własnego działania.

Według Encyklopedii PWN [www 1] sprzężenie zwrotne to „szczególny rodzaj oddziaływania między dwoma obiektami, odgrywający bardzo ważną rolę między innymi w automatyce, mechanice, elektronice, biologii, naukach społecznych”. Polega na zwrotnym oddziaływaniu skutku określonego zjawiska na jego przyczynę; według teorii sterowania zachodzi wtedy, gdy sygnał wyjściowy układu oddziałuje zwrotnie na jego sygnał wejściowy. Rysunek 4 przedstawia jeden z przykładów procesu zarządzania ryzykiem projektu – według metody PRINCE2, uwzględniający sprzężenia zwrotne.



**Rys. 4.** Proces zarządzania ryzykiem projektu w metodzie PRINCE2

Źródło: Shelton [2013, s. 6].

Schemat ten wyróżnia tylko dwie fazy: analizy ryzyka i zarządzania ryzykiem. Jednak wprowadzone sprzężenie zwrotne jest dość trywialne i nie uwzględnia wszystkich sytuacji, które mogą wystąpić podczas realizacji procesu zarządzania ryzykiem. Warto zadać pytanie: co się stanie na przykład w przypadku, gdy pojawią się ryzyka niezidentyfikowane w pierwszym kroku fazy

analizy ryzyka? Dlatego też autor proponuje uwzględnienie w procesie zarządzania ryzykiem projektu dwóch rodzajów sprzężeń zwrotnych:

- multisprzężenia zwrotnego,
- „samosprzężenia zwrotnego”.

Schemat blokowy na rys. 5 przedstawia koncepcję multisprzężenia zwrotnego.



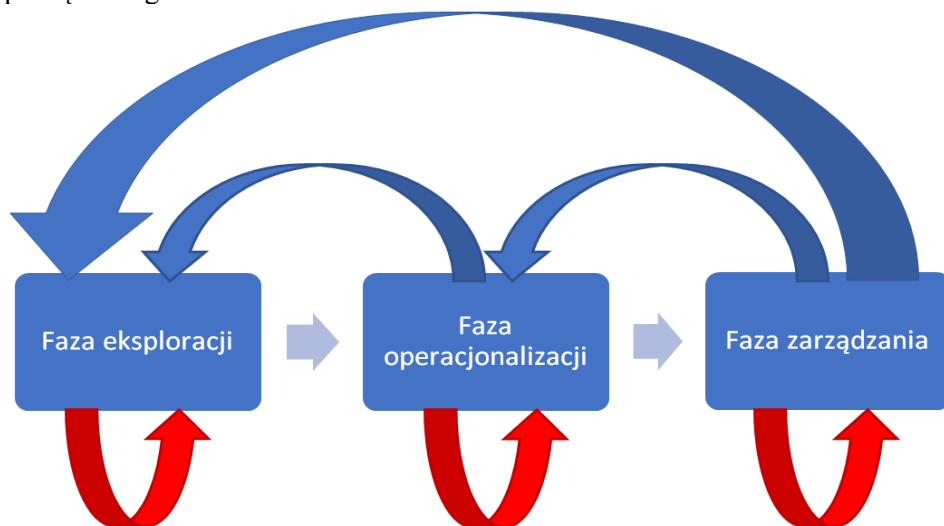
**Rys. 5.** Koncepcja multisprężenia zwrotnego w adaptacyjnym procesie zarządzania ryzykiem projektu – propozycja autora

W schemacie blokowym adaptacyjnej metody zarządzania ryzykiem sprzężenia zwrotne będą występować od każdego bloku do każdego innego bloku poprzedzającego. Przy czym sprzężenie takie biegnie nie tylko do bloku poprzedniego, ale do każdego bloku wcześniejszego. Jest zatem możliwy powrót z fazy zarządzania zarówno do fazy operacjonalizacji, jak i do fazy eksploracji ryzyka. Również z każdego kroku danej fazy można się wrócić do kroku z którejkolwiek fazy poprzedzającej. Z samej definicji sprzężenia zwrotnego: „doprowadzenie części sygnału wyjściowego z powrotem do wejścia” wynika, że takie przejście działa wyłącznie w tył, zatem nie jest możliwe bezpośrednie przejście z fazy eksploracji do fazy zarządzania.

Na rysunku 6 przedstawiono natomiast koncepcję „samosprężenia zwrotnego” (zaznaczoną strzałkami znajdującymi na dole każdego z bloków). Koncepcja ta polega na wprowadzeniu sprzężenia zwrotnego z danego bloku do tego samego bloku, w sposób iteracyjny, dopóki nie zostanie zrealizowana jedna z opcji:

- spełnienie danego warunku, co pozwoli na przejście do następnego bloku (podobnie jak w metodzie „stage-gate”),
- upływie czasu lub liczba iteracji, jakie zostały przewidziane dla danego warunku.

W takiej sytuacji – gdy warunek nie zostanie spełniony, a wystąpił timeout – może nastąpić: awaryjne przerwanie pętli i zakończenie procesu, powrót do bloku poprzedniego lub jednego z bloków poprzednich albo powrót do bloku początkowego.



**Rys. 6.** Koncepcja samosprzężenia zwrotnego w adaptacyjnym procesie zarządzania ryzykiem projektu – propozycja autora

Propozycje uwzględnienia obydwu przedstawionych sprzężeń zwrotnych zostały w tym rozdziale zarysowane w sposób ogólny. Ich pełna integracja z adaptacyjnym procesem zarządzania ryzykiem projektu będzie przedmiotem dalszych prac autora.

## Podsumowanie

W rozdziale omówiono proces zarządzania ryzykiem w projektach. Zaprezentowano propozycję trójfazowego schematu adaptacyjnego zarządzania ryzykiem w projektach, przedstawiono ideę sprzężeń zwrotnych w zarządzaniu ryzykami projektowymi oraz nakreślono koncepcję sprzężeń zwrotnych w procesie zarządzania ryzykiem w projekcie. Propozycja ta obejmuje wprowadzenie dwóch rodzajów sprzężeń zwrotnych: multisprzężenia zwrotnego oraz samo-

sprzężenia zwrotnego. Wprowadzenie zaproponowanych sprzężeń zwrotnych powoduje wzrost złożoności oraz czasochłonności procesu zarządzania ryzykiem związanych ze zwiększonym nakładem czasu i pracy w trakcie jego realizacji, co można uznać za słabą stronę przedstawionej propozycji. Do zalet tego rozwiązania można natomiast zaliczyć proaktywne podejście do realizacji procesu zarządzania ryzykiem, możliwość adaptacyjnej reakcji na pojawiające się ryzyka oraz zmienność ich atrybutów i parametrów.

Dalsze badania autora będą się koncentrowały nad rozwojem i dopracowaniem procesu adaptacyjnego zarządzania ryzykiem w projektach wytwórczych z branży urządzeń elektronicznych. Prace te w szczególności będą obejmowały: integrację zaproponowanych sprzężeń zwrotnych z adaptacyjną metodą zarządzania projektami wytwórczymi oraz przykłady praktycznego zastosowania w branży urządzeń elektronicznych.

## Literatura

- Allen J., Reichheld F., Hamilton B. (2005), *The Three "DS" of Customer Experience*, Harvard Management Update, November.
- Boehm B., Turner R. (2012), *Balancing Agility and Discipline*, Addison-Wesley, Boston.
- Jasińska K., Szapiro T. (2014), *Zarządzanie procesami realizacji projektów w sektorze ICT*, WN PWN, Warszawa.
- PMBOK Guide (2000), 2000 Edition, Project Management Institute, Newtown Square, Pennsylvania USA.
- PN-ISO 31000:2012 *Zarządzanie ryzykiem. Zasady i wytyczne*, <https://www.iso.org/pl/uslugi-zarzadzania/wdrazanie-systemow/zarzadzanie-ryzykiem/iso-31000/> (dostęp: 20.02.2021).
- Shelton D. (2013), *Risk Management and PRINCE2*, The Art of Service.
- Smith P.G. (2002), *Managing Risk Proactively in Product Development Projects*, IPLnet Workshop, 10-11.09.2002, Saas Fee, Switzerland.
- Smith P.G., Merritt G.M. (2002), *Proactive Risk Management*, Productivity Press, New York.
- Trocki M. (2017), *Metodyki i standardy zarządzania projektami*, PWE, Warszawa.
- Wyrozębski P. (2011), *Metodyka PRINCE2 [w:] Metodyki zarządzania projektami*, Bizarre, Warszawa.
- Wysocki R.K. (2013), *Efektywne zarządzanie projektami*, Helion, Gliwice.
- [www 1] <https://encyklopedia.pwn.pl> (dostęp: 20.02.2021).

## Propozycja globalnego wdrożenia systemu ERP z wykorzystaniem podejścia obiektowego

*Łukasz Tync<sup>1</sup>*

### Wprowadzenie

Znaczenie projektów i ich rozpowszechnienie zdecydowanie wzrasta w ostatnich dziesięcioleciach. Począwszy od opublikowanej w 1995 roku pracy pokazującej na przykładzie firmy Renault „projektyfikację” firm, aż po ostatnie lata, które wielu autorów wskazuje jako najdynamiczniejszy okres rozwoju zarządzania projektami. Przesuwa się też środek ciężkości z „prowadzenia większej ilości projektów” w kierunku „osiągania więcej dzięki projektom” [Maylor, Turkulainen, 2019]. Jest to szczególnie ważne dla organizacji intensywnie wykorzystujących wiedzę i projekty (KIPIO)<sup>2</sup> [Medina, Medina, 2015], ale globalizacja i cyfryzacja wpływają również na organizacje w tzw. tradycyjnych sektorach. Nawet jeśli firma działa w tradycyjnym biznesie, jej działy IT mogą w rzeczywistości realizować rozbudowane programy transformacji biznesowej. Różne działy organizacji mogą w różnym stopniu opierać się na wiedzy. Organizacje takie zostają narażone na ryzyko, złożoność i wyzwania, które wcześniej były ograniczone do firm działających w sektorach zaawansowanych technologii. Widoczne jest, że wiele firm konkuruje na rynkach, które podlegają gwałtownym zmianom z powodu wprowadzanych innowacji technologicznych. Z tego względu aktywnie poszukują one sposobu na orkiestrację swoich wysiłków, szybko zdając sobie sprawę, że traktowanie niektórych systemów jako niezależnych dziedzin jest w najlepszym przypadku nieoptymalne. Również patrząc z perspektywy samych projektów, coraz częściej źródła ryzyk i złożoności należy szukać poza nimi – czy to w czynnikach organizacyjnych, społeczno-ekonomicznych, czy technologicznych<sup>3</sup>. Nawet jeśli analizujemy perspektywę

<sup>1</sup> Uniwersytet Ekonomiczny w Katowicach, e-mail: lukasz.tync@sinapi.pl.

<sup>2</sup> Ang. Knowledge Intensive, Project Intensive Organisations.

<sup>3</sup> Znana firma doradcza Deloitte podkreśliła to w swoim corocznym raporcie na temat trendów technologicznych na 2018 rok, nadając mu tytuł „The symphonic enterprise” (Organizacja Symfoniczna). Zob. Deloitte [2018].

projektu, źródeł tych ryzyk i złożoności należy coraz częściej szukać poza samym projektem, w rosnącej dynamice zmian zachodzących w otoczeniu często opisywanych angielskim akronimem VUCA<sup>4</sup>.

W niniejszym rozdziale jako przykład programu transformacji zostanie wykorzystany globalny program wdrożenia systemu ERP (ang. Global ERP Rollout). Sam program będzie się składał z wielu podobnych projektów informatycznych oraz niejako wpływającej z nich transformacji procesów biznesowych realizowanych w zróżnicowanym środowisku organizacyjnym (spółki zależne tej samej organizacji korporacyjnej rozproszone geograficznie, posiadające wcześniej różne systemy i procesy, różną kulturę i regulacje prawne). Wraz ze wszystkimi otaczającymi projektami dotyczącymi tego samego systemu ERP program utworzy podportfolio<sup>5</sup> projektów organizacji.

Celem tego rozdziału jest budowa modelu otoczenia programu globalnego wdrożenia systemu ERP ze szczególnym uwzględnieniem interfejsów komunikacyjnych niezbędnych do prawidłowego osadzenia tego programu w organizacji. Praca jest częścią szerszych badań dotyczących wykorzystania pochodzącego z programowania podejścia obiektowego w zarządzaniu projektami, programami czy wręcz całym portfolio organizacji.

Autor niniejszego rozdziału zbuduje mapę zewnętrznych zależności, a także zidentyfikuje interfejsy, które powinny pomóc zmniejszyć ryzyka programu pochodzące z jego otoczenia i zwiększyć wpływ programu poprzez usprawnienie transferu wiedzy oraz zarządzanie procesem uczenia się organizacji i zarządzania wiedzą.

Przykładowy projekt będzie fuzją rzeczywistych doświadczeń, które autor zebrał w ciągu wielu lat pracy jako menedżer i konsultant, uczestnicząc w podobnych programach. Pomimo faktu, że implementacja będzie specyficzna dla danego projektu, można założyć, że proces i korzyści płynące z budowy modelu mogą być powielane również w innych typach organizacji.

---

<sup>4</sup> Ang. Volatility – zmienność, Uncertainty – niepewność, Complexity – złożoność, Ambiguity – dwuznaczność.

<sup>5</sup> W zależności od wersji angielska wersja PMBoK odwołuje się do „subportfolio” lub „subsidiary portfolio”.

### 3.1. Przegląd literatury

Już w 1998 roku Jaafari i Manivong [1998] postulowali stworzenie nowych, „inteligentnych” systemów informatycznych do zarządzania projektami jako odpowiedź na rosnącą złożoność i niepewność projektów i ich otoczenia. Podkreślali również potrzebę stworzenia struktury komunikacji i integracji działań projektowych opartej na modelu informacyjnym unikalnym dla każdego projektu. Z kolei Rolstadås i Schiefloe [2017] zwracali uwagę na wielu autorów podkreślających, że niezdolność do zarządzania rosnącą złożonością projektów jest powszechnie uznawana za główny powód, dla którego wiele projektów nie osiąga swoich celów [Brady, Davies, 2014]. Krytykowali oni dominujący paradygmat w zarządzaniu projektami za brak spójności i poświęcanie zbyt małej uwagi kontekstowi zarządzaczemu, w szczególności złożoności projektu i jego otoczenia [Cicmil i in., 2009], twierdząc, że publikacje na tym polu mają charakter wysoce nakazowy i ignorują kontekst projektu.

Biorąc pod uwagę duży odsetek projektów zakończonych porażką, zarządzający przedsiębiorstwami zaczynają kwestionować rzeczywistą wartość dostarczaną przez projekty, programy i portfele projektów. Z tego względu należy położyć nacisk na odpowiednio wdrożony ład korporacyjny<sup>6</sup> (ang. Governance), dzięki któremu organizacja może zapewnić sobie wykonalność projektów i odpowiednio ich ukierunkowanie [Müller, Lecoeuvre, 2014]. Organizacje rozwijają ład korporacyjny w obszarze zarządzania projektami, aby upewnić się, że projekty i programy generują wartość dla organizacji bazowej [Joslin, Müller, 2016]. Jest to osiągnięte poprzez tworzenie i pielęgnowanie powiązań wewnątrz i pomiędzy różnymi domenami w organizacji. Riis, Hellström i Wikström [2019] argumentują, że to właśnie powiązania organizacyjne, a nie struktura organizacyjna w konwencjonalnym sensie, mają kluczowe znaczenie dla generowania wartości. Autorzy doszli do wniosku, że projekty nie powinny być traktowane jak oddzielne wyspy [Engwall, 2003], a raczej jako elementy głęboko osadzone w stałej organizacji, z wieloma powiązaniami z ich kontekstem organizacyjnym [Sydow, Lindkvist, Defillippi, 2004].

Brady i Davies [2014] definiują dwa poziomy uczenia się w organizacjach opartych na projektach (ang. Project Based Organisations – PBO): uczenie prowadzone przez projekty (uczenie się oddolne – od projektów do organizacji ma-

---

<sup>6</sup> Alternatywnym tłumaczeniem angielskiego słowa „governance” byłoby zarządzanie, jednakże w takim wypadku pominięte zostałyby istotne rozróżnienie pomiędzy „project management” a „governance of projects”.

cierzystej) oraz uczenie, któremu przewodzi biznes (przenoszenie wiedzy z organizacji macierzystej do projektów). Jednym z założeń tego rozdziału jest stworzenie efektywnej pętli wiedzy, która powinna umożliwić właściwą koordynację pomiędzy projektem a organizacją nadrzędną.

### **3.2. Podejście obiektowe w zarządzaniu programem**

Rosnąca złożoność i niepewność, przed którą stają zarządzający projektami, programami czy portfelami projektów w organizacji, wydaje się mieć wiele analogii do wyzwania, przed którym w przeszłości stawali programiści i projektanci oprogramowania, dla których remedium okazało się podejście obiektowe. W swoich wcześniejszych publikacjach [Tync, 2012a, 2012b] autor podejmował już próby przeniesienia elementów podejścia obiektowego na grunt zarządzania projektami, ale temat ten wciąż wymaga rozwinięcia i usystematyzowania.

Autor postuluje zdefiniowanie programu jako sieci luźno powiązanych obiektów. Obiekty będą zazwyczaj reprezentować określony rodzaj odpowiedzialności (rolę lub zadanie do wykonania) i powinny być możliwie autonomiczne. Sieć luźno powiązanych obiektów powinna pomóc nam zbudować elastyczny model portfela projektów, ale nie będzie to kompletne, dopóki nie przeniesiemy go na pewien abstrakcyjny poziom. Podążając za analogią programistyczną, można powiedzieć, że klasy będą stabilnymi abstrakcjami używanymi głównie w fazie planowania, podczas gdy obiekty będą jej instancjami w czasie rzeczywistym. Bardzo często możemy mieć wiele instancji tej samej klasy (np. wiele różnych zadań i wielu programistów te zadania realizujących).

Założmy, że projekt jest pewną częścią rzeczywistości, którą chcemy kontrolować. Plan projektu w takim przypadku byłby jego uproszczonym modelem. Oczekiwany model powinien dawać menedżerom możliwość zrozumienia i opisanie relacji pomiędzy różnymi zadaniami (np. zadanie A jest poprzednikiem zadania B) i zasobami (do wykonania zadania C potrzebny jest zasób A). Powinien również dawać kierownikom projektów możliwość monitorowania postępów i porównywania rzeczywistej sytuacji z poziomem bazowym. W idealnym przypadku powinien również dać im narzędzia do reagowania, jeśli dane rzeczywiste za bardzo odbiegają od oczekiwań. Model taki może też objąć kluczowe otoczenie projektu oraz brać pod uwagę jego specyfikę organizacyjną, żeby pomóc zarządzającym zidentyfikować i kontrolować istotne połączenia pomiędzy projektem a jego otoczeniem.

### 3.3. Model otoczenia programu

#### 3.3.1. Założenia

W niniejszym przykładzie posłużono się modelową firmą, która w istotnym stopniu nie odbiega od rzeczywistych organizacji, z którymi autor współpracował w przeszłości. Jest to tradycyjna organizacja przemysłowa, która ze względu na wykorzystanie powszechnych w biznesie systemów informatycznych oraz globalizację jest wystawiana w ostatnich latach na dużo większą zmienność i złożoność. Poniższa charakterystyka z jednej strony odpowiada wielu istniejącym na rynku firmom, z drugiej jest na tyle dokładna, że dobrze ilustruje wyzwania stojące przed taką organizacją:

- ponad dziesięć tysięcy pracowników na całym świecie,
- spółki zależne zlokalizowane na całym świecie: co najmniej jedna fabryka w Europie, Azji, Stanach Zjednoczonych Ameryki i Ameryce Południowej oraz sieć oddziałów handlowych w najważniejszych krajach,
- działalność w „tradycyjnym” przemyśle, np. inżynierskim lub infrastrukturalnym,
- długa historia obejmująca fuzje i przejęcia firm różnej wielkości,
- heterogeniczność kulturowa i rozproszone środowisko decyzyjne,
- szeroko wdrożona centralizacja i wirtualizacja infrastruktury IT (na poziomie platformy),
- szeroki przekrój starszych systemów ERP pochodzących od różnych dostawców, zazwyczaj starych i niezależnych (warstwa aplikacji),
- szereg dodatkowych adaptacji i specjalistycznych aplikacji zintegrowanych z istniejącymi systemami ERP w celu pokrycia luk funkcjonalnych,
- rozproszona organizacja IT z widocznymi silosowymi aplikacjami w większości przypadków fragmentarycznie rozwiązującymi określone problemy biznesowe i połączonymi poprzez szereg interfejsów,
- strategiczna wizja wdrożenia scentralizowanego dwupoziomowego podejścia do ERP – SAP wybrany i wdrażany jako podstawowy system ERP dla fabryk i regionalnych centrów dystrybucji oraz iScala<sup>7</sup> jako docelowe rozwiązanie dla lokalnych oddziałów sprzedaży; oba ERP mają być wdrożone na podstawie ustandaryzowanych szablonów znanych zwykle jako Globalny Szablon ERP (ang. Global ERP Template),

---

<sup>7</sup> Oba wspomniane systemy (SAP i iScala) są istniejącymi na rynku od wielu lat rozwiązaniami klasy ERP. SAP jest dominującym wyborem w dużych organizacjach przede wszystkim w Europie, podczas gdy iScala jest wybierana przez średnie firmy bądź przez firmy duże jako drugi poziom ERP (dla mniejszych i mniej złożonych oddziałów).

- organizacja ma standardową metodykę zarządzania projektami, która jest w rzeczywistości ogólnym zestawem wytycznych i szablonów dokumentów, niewystarczająco szczegółowym, aby kompleksowo przeprowadzić globalne wdrożenie szablonu ERP,
- na ogół organizacja realizuje projekty IT, korzystając z wewnętrznych zasobów, ale w przypadku większych projektów (takich jak wdrożenia ERP) musi korzystać również z zasobów zewnętrznych,
- w niektórych krajach wprowadzono korporacyjne centrum usług wspólnych<sup>8</sup> dla obszaru finansów, które jest systematycznie rozwijane.

Rozpatrywany przykład będzie skupiony na wdrożeniu podejścia obiektowego z perspektywy biura zarządzania projektami (PMO)<sup>9</sup> na poziomie drugiego systemu ERP. Powyższe założenia ilustrują poziom i potencjalne źródła złożoności, jakie mogą napotkać zarządzający programem. Program składa się ze stosunkowo prostych, powtarzalnych projektów, które będą jednak realizowane w heterogenicznym i bardzo złożonym otoczeniu. Możemy zaobserwować zarówno różnorodność kulturową, jak i wynikającą z różnych systemów prawnych czy zastanych rozwiązań technicznych (infrastruktura IT) bądź procesowych. Dodatkowym źródłem niepewności jest także duża i niejednorodna grupa interesariuszy każdego z projektów.

Założono, że wstępny szablon (global template) jest gotowy do wdrożenia, ale pierwsze projekty będą traktowane jako wdrożenia pilotażowe i mogą mieć istotny wpływ na ostateczny kształt szablonu. Wstępny zakres jest ograniczony do 21 wdrożeń, przy czym maksymalnie trzy projekty będą prowadzone równolegle. Pojedyncze wdrożenie będzie trwało od 4 do 8 miesięcy (średnio 6), zatem przewidywany czas realizacji to 3,5 roku do 4 lat. Po przeprowadzeniu projektów pilotażowych zostanie zorganizowany warsztat podsumowujący, którego celem będzie zebranie i usystematyzowanie doświadczeń oraz zapewnienie, że zebrana wiedza będzie dostępna dla przyszłych wdrożeń. Ze względu na długą kolejkę spółek oczekujących na wdrożenie zarząd może podjąć decyzję o zwiększeniu liczby równolegle prowadzonych wdrożeń lub wywierać nacisk na skrócenie średniego czasu trwania projektu.

---

<sup>8</sup> Według Bergerona [2003] centrum usług wspólnych (ang. shared service center) jest strategią współpracy, w której podzbiór istniejących funkcji biznesowych jest skoncentrowany w nowej, pół-autonomicznej jednostce biznesowej, która ma strukturę zarządzania zaprojektowaną w celu promowania wydajności, generowania wartości, oszczędności kosztów i poprawy obsługi wewnętrznych klientów korporacji macierzystej, jak w przypadku firmy konkurującej na otwartym rynku.

<sup>9</sup> Ang. Project Management Office – dedykowana jednostka wspierająca kierowników projektów i dbająca o zachowanie standardów korporacyjnych w projektach objętych jego zasięgiem odpowiedzialności.

Ponieważ nie ma jeszcze dedykowanego rozwiązania informatycznego wspierającego prezentowane podejście, a pierwsze próby nie powinny się wiązać z dużymi inwestycjami, założono maksymalne wykorzystanie narzędzi już istniejących w organizacji. Mimo że przykład dotyczy dużego przedsiębiorstwa działającego w „tradycyjnej” branży, omawiane projekty będą prowadzone przez dział IT, który jest wewnętrzną organizacją informatyczną działającą zgodnie z branżowymi standardami i wykorzystującą najpopularniejsze systemy ułatwiające zarządzanie projektami informatycznymi i operacjami. Na tej podstawie można oczekiwać co najmniej podstawowego wdrożenia standardu ITIL<sup>10</sup> z Katalogiem Usług i systemem zgłoszeniowym po stronie operacyjnej (wybrano OTRS<sup>11</sup> jako bardzo popularne rozwiązanie open source) oraz JIRA<sup>12</sup> do zarządzania rozwojem oprogramowania i adaptacjami. Autor zakłada, że istnieje również scentralizowane oprogramowanie do zarządzania dokumentami. Nowoczesne systemy informatyczne są projektowane tak, aby były elastyczne i pozwalały nie tylko na dostosowanie do własnych potrzeb (coraz częściej dostępne dla zaawansowanych użytkowników końcowych bez umiejętności programistycznych), ale także na standardowe sposoby integracji i przekazywania danych (np. za pomocą ustandaryzowanych, otwartych API<sup>13</sup>). Używane w organizacji narzędzia zostaną zaadoptowane, a ewentualne braki organizacyjne zostaną uzupełnione „papierowymi”<sup>14</sup> procedurami, które będą realizowane przez PMO lub kierownika projektu.

### 3.3.2. Obszar modelu

Jak już wspomniano, na podstawie decyzji korporacyjnej został utworzony dział PMO dla programu wdrożenia ERP drugiego poziomu. Uznano, że obszar zarządzany w tym przykładzie zostanie zdefiniowany jako wszystkie projekty i ope-

<sup>10</sup> ITIL (ang. IT Infrastructure Library) jest publicznym zbiorem procesów i najlepszych praktyk w zarządzaniu usługami IT. Dostarcza on ram dla zarządzania IT i skupia się na ciągłym pomiarze i poprawie jakości dostarczanych usług IT, zarówno z perspektywy biznesu, jak i klienta. ITIL jest powszechnie stosowany przez organizacje wdrażające techniki i procesy w swoich organizacjach [Cartlidge i in., 2007].

<sup>11</sup> Początkowo funkcjonujący pod nazwą angielską Open-Source Ticket Request System (OTRS) – pakiet narzędzi wspierających zarządzanie usługami IT.

<sup>12</sup> Autorski produkt firmy Atlassian umożliwiający śledzenie błędów i zwinne zarządzanie projektami.

<sup>13</sup> Ang. Application Programming Interface – interfejs programowania aplikacji zapewnia programistyczny interfejs do komponentu oprogramowania lub usługi. Pozwala on programistom, innym niż ci, którzy napisali dany komponent, na korzystanie z określonego komponentu lub usługi.

<sup>14</sup> W tym kontekście za „papierowe” będą uznane procedury, które nie są wspierane bezpośrednio przez dedykowane rozwiązanie IT, jednak wciąż mogą być realizowane z użyciem poczty email bądź cyfrowego przypomnienia w aplikacji kalendarza itp.

racje związane z tym konkretnym środowiskiem ERP (w większości przypadków na jednym serwerze aplikacyjnym może działać wiele podmiotów). Jest to obszar szerszy od zakresu odpowiedzialności samego PMO, który ogranicza się do programu wdrożenia ERP, a nie bierze pod uwagę operacji związanych z utrzymaniem już wdrożonych jednostek czy też zmian, które mogą z tych operacji wynikać.

Ta różnica może nie być widoczna w początkowej fazie programu (kiedy najprawdopodobniej użyty będzie nowy zestaw serwerów czekających na pierwsze uruchomienia), ale jej znaczenie będzie rosło wraz z rosnącą liczbą oddziałów działających produkcyjnie w systemie. Kiedy pierwsze filie zaczną korzystać z systemu w sposób produkcyjny, nie będą czekać, aż wszystkie pozostałe projekty zostaną zakończone. Zaczną zgłaszać wnioski o zmiany, aby podążać za trendami rynkowymi. Niektóre z tych zmian będą adekwatne tylko lokalnie, ale niektóre mogą się stać na tyle istotne, że organizacja będzie zmuszona włączyć je do globalnego szablonu. Najprawdopodobniej będzie musiała się również dostosować do nowych przepisów prawnych.

W przypadku wątpliwości warto poszukać elementów, które wnoszą istotną złożoność i starać się unikać dzielenia takich elementów<sup>15</sup>. W omawianym przypadku takim elementem złożoności jest samo środowisko ERP. Gdyby zdecydowano się ograniczyć zarządzany obszar tylko do programu, to i tak należałoby zaprojektować skomplikowany mechanizm zarządzania ograniczeniami i synchronizacji pomiędzy toczącymi się projektami wdrożeniowymi a operacjami w działających już produkcyjnie podmiotach. Rozszerzając ten obszar, uzyskuje się możliwość modelowania i kontroli tych relacji w ramach podejścia obiektowego<sup>16</sup>.

### 3.3.3. Identyfikacja klas

Mimo że niniejszy rozdział koncentruje się na otoczeniu programu, a nie na samym programie, należy rozpocząć od zidentyfikowania kluczowych klas wymaganych do zbudowania całego modelu. Ten krok obejmie zarówno model programu, jak i jego otoczenie, ale tutaj skupiono się w szczególności na tym drugim.

---

<sup>15</sup> Sytuacja, w której część elementu znalazłaby się wewnątrz obszaru, a część poza nim, byłaby trudna do zarządzania i wymagałaby dużych nakładów pracy związanych z koordynacją i negocjacjami stanowisk interesariuszy.

<sup>16</sup> Oznacza to, że modelowany obszar będzie obejmował program globalnego wdrożenia ERP rozszerzony o operacje i przyszłe projekty dotyczące tego samego środowiska ERP. Dla uproszczenia w rozdziale będzie używane określenie „program”, mimo że z czasem bardziej adekwatne mogłoby być określenie „subportfolio” projektów.

Tabela 1 przedstawia klasy zidentyfikowane dla omawianego przykładu. Biorąc pod uwagę, że jest to wdrożenie oprogramowania i jednocześnie program transformacji biznesowej realizowany przez wysoko wykwalifikowanych specjalistów, szczególnej uwagi będzie wymagało zarządzanie zasobami ludzkimi. Wymagane będą również znacząca koordynacja zaangażowanych jednostek organizacyjnych (rozproszonych globalnie) oraz właściwe zarządzanie infrastrukturą IT.

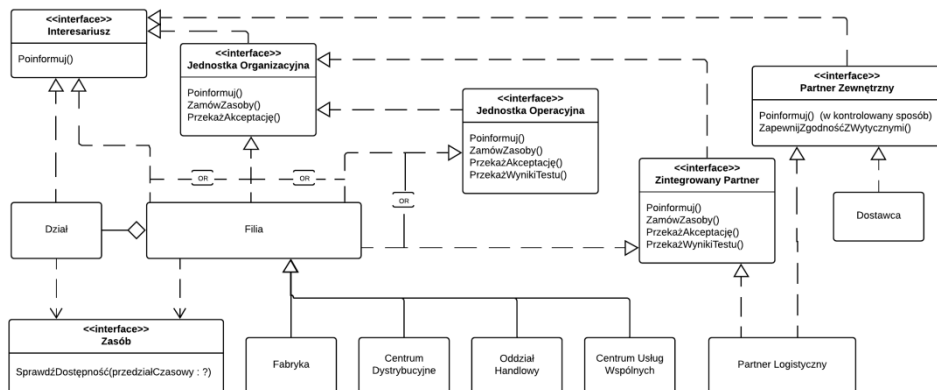
**Tabela 1.** Identyfikacja klas dla przykładowego programu globalnego wdrożenia systemu ERP

| Zewnętrzne/ połączenia                                                                                                                                                                                                                                                                                      | Program                                                                                                                                                                                                                                               |                                                                                                                                                                               | Niezdefiniowane                                                                                                                 |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------|
| <b>Organizacja</b>                                                                                                                                                                                                                                                                                          | <b>Zarządzanie projektami/ programami</b>                                                                                                                                                                                                             |                                                                                                                                                                               | <b>Produkt/ rezultat projektu</b>                                                                                               |
| <ul style="list-style-type: none"> <li>• Dział</li> <li>• Jednostka organizacyjna</li> <li>• Interesariusz</li> <li>• Filia</li> <li>• Fabryka</li> <li>• Centrum dystrybucyjne</li> <li>• Oddział handlowy</li> <li>• Centrum usług wspólnych</li> <li>• Zintegrowani partnerzy</li> </ul>                 | <ul style="list-style-type: none"> <li>• Projekt</li> <li>• Program</li> <li>• Portfel projektów</li> <li>• Czynność/ zadanie</li> <li>• Plan projektu</li> </ul>                                                                                     | <ul style="list-style-type: none"> <li>• Status</li> <li>• Budżet</li> <li>• Jakość</li> <li>• Ryzyko</li> <li>• Cel</li> <li>• Biuro zarządzania projektami (PMO)</li> </ul> | <ul style="list-style-type: none"> <li>• Dokument</li> <li>• Zakres</li> <li>• Backlog</li> <li>• Wymaganie</li> </ul>          |
| <b>Infrastruktura</b>                                                                                                                                                                                                                                                                                       | <ul style="list-style-type: none"> <li>• Plan programu</li> <li>• Project Charter</li> <li>• Change Request</li> <li>• Artifact Templates</li> </ul>                                                                                                  | <ul style="list-style-type: none"> <li>• Governance</li> <li>• Metodologia (ogólne wytyczne)</li> </ul>                                                                       |                                                                                                                                 |
| <b>Wiedza i procesy</b>                                                                                                                                                                                                                                                                                     | <b>Wymagające uszczegółowienia</b>                                                                                                                                                                                                                    |                                                                                                                                                                               | <b>„Parking”</b>                                                                                                                |
| <ul style="list-style-type: none"> <li>• Proces</li> <li>• Dokument</li> <li>• Zarządzanie zmianą</li> <li>• Właściciel procesu</li> <li>• Wiedza</li> <li>• Cele strategiczne i pomiary (KPIs)</li> <li>• Kompetencje</li> <li>• Informacja biznesowa</li> <li>• Procedury i instrukcje robocze</li> </ul> | Zarządzanie zasobami (ludzkimi): <ul style="list-style-type: none"> <li>• Zasoby</li> <li>• Dostawca zasobów</li> <li>• Zespół</li> <li>• Zasoby wewnętrzne/ zewnętrzne</li> <li>• Zasoby ludzkie</li> <li>• Sprzęt</li> <li>• Kompetencje</li> </ul> |                                                                                                                                                                               | <ul style="list-style-type: none"> <li>• Regulacje</li> <li>• Baza wiedzy</li> <li>• Kamień milowy</li> <li>• Sukces</li> </ul> |

Źródło: Opracowanie własne.

### 3.3.4. Organizacja

W tym obszarze zebrano pojęcia związane z potencjalnymi klasami wymaganymi w modelu. Graficzna reprezentacja modelu została przygotowana z wykorzystaniem elementów notacji UML<sup>17</sup>. Głównym kryterium wyboru klas jest ich potencjalna rola lub związek z projektem. Przedstawiony na rys. 1 model organizacji pozwala również określić kluczowe interfejsy, które będą potrzebne.



**Rys. 1.** Model obszaru organizacji dla przykładu globalnego wdrożenia systemu ERP

Źródło: Opracowanie własne.

Jednym z założeń powyższego przykładu było to, że organizacja referencyjna posiada wiele filii (odrębnych podmiotów prawnych) rozproszonych globalnie i pełniących różne role (fabryka, centrum dystrybucyjne, oddział handlowy, centrum usług wspólnych itp.). Każda z filii może być podzielona na działy i zarówno filia, jak i konkretny dział może potencjalnie odgrywać rolę interesariusza projektu lub posiadać zasoby niezbędne do osiągnięcia celów projektu. Zidentyfikowano pięć głównych interfejsów tworzących hierarchiczną strukturę (interesariusz, jednostka organizacyjna, jednostka operacyjna, zintegrowany partner, partner zewnętrzny). W zależności od roli i powiązań z programem spółka zależna będzie implementować jeden z nich. Najbardziej ogólnym interfejsem jest Interesariusz (dla uproszczenia ograniczono kluczowe metody tylko do Poinformuj()). Dla bardziej zaangażowanych jednostek organizacyjnych (np. Centrum Usług Wspólnych, które może potencjalnie uczestniczyć w wielu projektach poprzez udostępnianie zasobów i zatwierdzanie pewnych procesów)

<sup>17</sup> Ang. Unified Modelling Language.

stworzono interfejs Jednostka Organizacyjna. Pozostałe dwa – na obecnym etapie – wydają się mieć dokładnie takie same metody, ale ze względu na różne role można założyć, że w przyszłości będą się one różnić. Jednostka Operacyjna (w niektórych firmach jest używany skrót OPCO od angielskiego „operational company”) będzie jednostką, która jest właścicielem pojedynczego projektu i jego głównym beneficjentem, natomiast Zintegrowany Partner będzie innym podmiotem prawnym połączonym z firmą operacyjną poprzez przynajmniej jeden łańcuch procesów biznesowych, zazwyczaj realizowanych z wykorzystaniem mechanizmu EDI<sup>18</sup>. W niektórych przypadkach rola partnera zintegrowanego może być zlecona firmie zewnętrznej, dlatego zdefiniowano również interfejs Partner Zewnętrzny, który może wymagać większej kontroli ze względu na interakcję wychodzącą poza obszar korporacji.

### 3.3.5. Infrastruktura

W prezentowanym przykładzie skupiono się głównie na infrastrukturze IT, dzięki czemu możliwe jest wykorzystanie istniejących narzędzi i metod pochodzące ze standardów branżowych. Jednym z kluczowych procesów ITIL jest Service Asset and Configuration Management (SACM), którego celem jest identyfikacja, kontrola i rozliczanie aktywów usługowych i elementów konfiguracji (CI)<sup>19</sup>, jak również ochrona i zapewnienie ich integralności w całym cyklu życia usługi. W dużych organizacjach proces ten jest wspierany przez bazę danych zarządzania konfiguracją (CMDB)<sup>20</sup>.

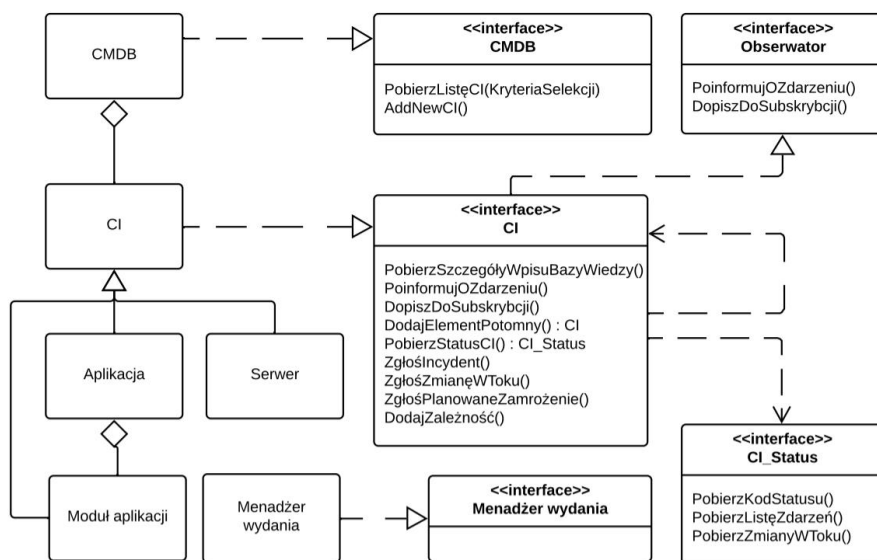
Zarówno CI, jak i CMDB zostaną bezpośrednio wykorzystane w części modelu dotyczącej infrastruktury, aby zapewnić odpowiednie dopasowanie wszystkich działań związanych z szeroko rozumianą infrastrukturą IT (w tym określonych komponentów lub modułów systemu ERP), ale także jako część procesu zarządzania wiedzą. Umożliwi to również bezpośrednie połączenie z rezultatem każdego projektu – skonfigurowanym i dostosowanym do potrzeb systemem oraz powiązanymi z nim procesami biznesowymi. Rysunek 2 przed-

<sup>18</sup> EDI to skrót od angielskiego Electronic Data Interchange – elektroniczna wymiana danych.

<sup>19</sup> Ang. Configuration Items to komponenty infrastruktury informatycznej (lub elementy związane z tą infrastrukturą), które podlegają kontroli w ramach procesu Zarządzania Konfiguracją.

<sup>20</sup> Ang. Configuration Management Database jest logicznym modelem infrastruktury i usług IT, którego tworzenie i utrzymywanie jest głównym rezultatem procesu Zarządzania Konfiguracją; równolegle jest ona centrum informacyjnym gromadzącym informacje o stanie poszczególnych elementów konfiguracji (CI).

stawia wstępny model infrastruktury oparty na trzech kluczowych interfejsach – CMDB, CI oraz Status\_CI<sup>21</sup>. Założono, że CMDB będzie zawierać hierarchiczną strukturę elementów konfiguracji, które mogą wchodzić w relacje (jeden rodzic może zawierać wiele elementów potomnych) i posiadać zależności. Istotne okazało się również utrzymanie informacji o trwających incydentach, zmianach<sup>22</sup> lub specjalnych okresach „zamrożenia”<sup>23</sup>, kiedy jakkolwiek zmiana w danym CI nie powinna być dozwolona. Założono również, że korzystne będzie, jeśli CI będą dziedziczyły po istniejącym interfejsie Obserwator. Prawdopodobnie potrzebny będzie również interfejs dla klasy, która będzie pełniła rolę Menadżera Wydania<sup>24</sup>, ale na ten moment nie zidentyfikowano jeszcze wyraźnego powiązania z modelem.



**Rys. 2.** Model infrastruktury dla przykładu globalnego wdrożenia systemu ERP

Źródło: Opracowanie własne.

<sup>21</sup> Zawierający ustrukturalizowaną informację o statusie poszczególnych elementów konfiguracji (np. czy element działa prawidłowo itp.).

<sup>22</sup> Incydenty i zmiany są specyficznymi pojęciami zaczerpniętymi z ITIL i dotyczą statusu poszczególnych elementów konfiguracji.

<sup>23</sup> Ang. Freeze Period – okres, w którym zabronione jest wprowadzanie zmian w konfiguracji (np. ze względu na krytyczne działania biznesowe, takie jak zamknięcie miesiąca w dziale księgowości lub wyprzedaże świąteczne dla działu sprzedaży).

<sup>24</sup> Ang. Release Manager.

### **3.3.6. Wiedza i procesy**

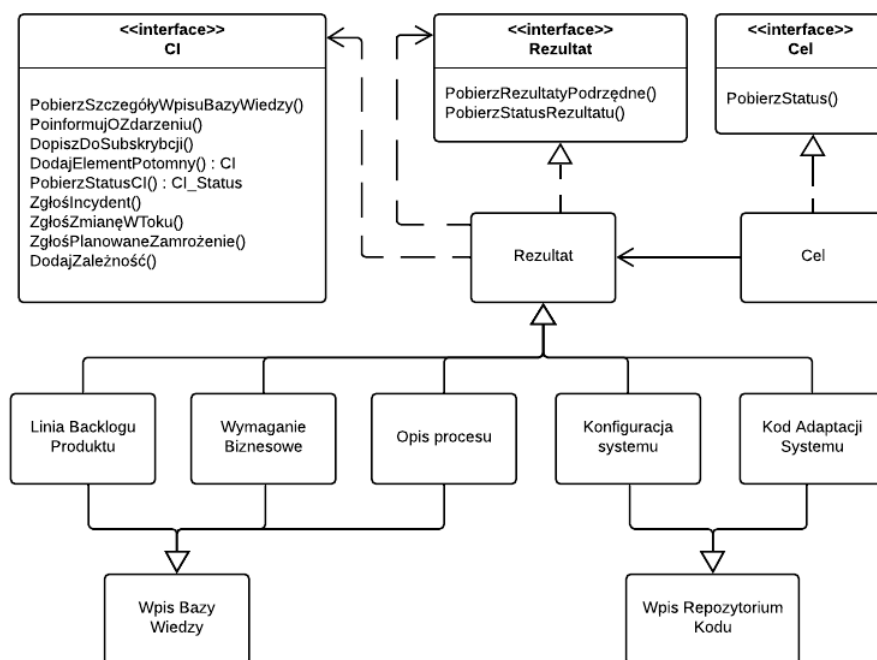
Obszar wiedza i procesy został zachowany na koniec ze względu na jego rolę łącznika zarówno pomiędzy projektem a jego otoczeniem, jak i pomiędzy klasami zidentyfikowanymi już w zewnętrznych obszarach. Jednym z celów w tym przypadku będzie zapewnienie efektywnej pętli wiedzy, która powinna umożliwić właściwą koordynację pomiędzy organizacją prowadzącą projekt a organizacją macierzystą. Koordynacja ta powinna zapewnić, że istniejąca wiedza, procesy i cele strategiczne są kaskadowane do projektów, a wiedza zgromadzona w ramach projektów – w tym zmodyfikowane procesy – powinna być transparentnie przekazywana do organizacji macierzystej. Strategię korporacyjną zdekomponowano na cele strategiczne i KPI<sup>25</sup>, które powinny być z kolei przeniesione na cele projektu lub programu i dalej na cele poszczególnych zadań. Strategia korporacyjna będzie pośrednio związana z procesami biznesowymi, które będą przekształcane poprzez projekty (z wykorzystaniem zarządzania zmianą). Dokumenty projektowe procesów staną się częścią bazy wiedzy, która powinna zawierać nie tylko inne rodzaje artefaktów projektowych, ale również odniesienia do CMDB, która jest ważnym elementem zarządzania wiedzą. Model zaprezentowano na rys. 3.

---

<sup>25</sup> Ang. Key Performance Indicator – mierzalna wartość ilustrująca, w jakim stopniu organizacja lub jednostka realizuje określony cel biznesowy.



biznesowych, które mogą budować backlog produktu, jak i procesy biznesowe objęte tym wdrożeniem. Można też zdekomponować ten rezultat na konfigurację systemu, kod programu (w przypadku adaptacji), podręczniki użytkownika i inne dokumenty procesowe itp. W niektórych przypadkach można uzyskać bezpośredni dostęp do cyfrowego produktu (np. kodu programu dostępnego w systemach dostarczania i kontroli wersji, jak Bitbucket<sup>26</sup>), w innych będzie istnieć jego cyfrowa reprezentacja (dokumenty projektowe procesów lub podręczniki). Niezbędny jest jednak uniwersalny sposób (interfejs) dostępu do tych obiektów oraz zapewnienie ich zgodności, spójności i jednoznaczności. Jak widać na rys. 4, interfejs ten powinien mieć bezpośrednie połączenie z wcześniej zdefiniowaną CMDB (poprzez interfejs CI) i będzie albo wpisem do Bazy Wiedzy, albo elementem repozytorium kodu. Ze względu na swój charakter może być również wykorzystywany przez klasę Cel.



**Rys. 4.** Model rezultatu dla przykładu globalnego wdrożenia systemu ERP

Źródło: Opracowanie własne.

<sup>26</sup> Bitbucket to autorskie oprogramowanie repozytorium kodu oferowane przez firmę Atlassian zarówno w modelu Software as a Service, jak i On-premise (zainstalowane indywidualnie przez klienta na własnym serwerze).

### 3.3.8. Pętla zarządzania wiedzą

Opisane powyżej obszary nakładają się na siebie, a niektóre klasy mogą odgrywać rolę w wielu z nich. W dalszej części (rys. 4-5) zostaną połączone uproszczone modele wszystkich obszarów, aby zbudować kompletny model otoczenia programu i odwzorować istotne punkty styku, które wymagają odpowiedniego uwzględnienia w budowanym modelu programu. Dla pełniejszego obrazu zostaną dodane dwa dodatkowe interfejsy (Menadżer Zmiany i Menadżer Zasobów).

Jeśli połączy się zdefiniowane klasy w grupy na podstawie pierwotnie zdefiniowanych obszarów zewnętrznych, okaże się, że model odzwierciedla wspomnianą już pętlę kompetencji na obu poziomach uczenia (rys. 4-6). Widoczne na rysunku strzałki wskazują z jednej strony kierunek przepływu istniejącej wiedzy, procesów i celów strategicznych z organizacji macierzystej do projektów, a z drugiej przepływ wiedzy zgromadzonej w ramach projektów, w tym wiedzy o zmodyfikowanych procesach, w kierunku organizacji macierzystej.

Zamodelowano kluczowe zależności zewnętrzne, które powinny być dalej ustrukturalizowane i zarządzane przy użyciu standardowych interfejsów. Będzie to istotne dla zmniejszenia ryzyka pochodzącego z otoczenia programu i zwiększenia jego wpływu poprzez usprawnienie transferu wiedzy oraz ułatwienie procesu uczenia się organizacji i zarządzania wiedzą.





## Podsumowanie

W niniejszym rozdziale podjęto próbę budowy obiektowego modelu otoczenia organizacji projektowej oraz zdefiniowania istotnych punktów styku pomiędzy projektem a jego otoczeniem. Na tej podstawie możliwe jest zdefiniowanie kluczowych interfejsów strukturalizujących obustronny przepływ informacji pomiędzy uczestnikami projektu a organizacją zewnętrzną dla zapewnienia prawidłowej pętli sprzężenia zwrotnego. Pętla taka jest niezwykle istotna w organizacjach intensywnie opartych na wiedzy i projektach. Proponowane podejście z jednej strony wydaje się uniwersalne i możliwe do zastosowania w różnych branżach i organizacjach, z drugiej strony pozwala osobie budującej model w bardzo precyzyjny sposób uchwycić specyfikę organizacji.

W idealnej sytuacji organizacja miałaby do dyspozycji dedykowany system wspomagający zarządzanie projektami implementujący podejście obiektowe. Mimo że rozważana jest budowa takiego systemu jako jeden z potencjalnych następnych kroków, wykracza to daleko poza zakres tego rozdziału ze względu na złożoność i koszty. Proponowane podejście może przynieść wartość dodaną nawet bez dedykowanego systemu dzięki maksymalnemu wykorzystaniu istniejących zwykle w organizacji narzędzi. Autor proponuje adaptację istniejącej technologii i wypełnienie ewentualnych luk organizacyjnych „papierowymi”<sup>27</sup> procedurami, które będą wykonywane przez PMO lub kierownika projektu. Rozpoczęto analizę od modelu, który w kolejnych krokach należałoby odwzorować w narzędziach i procedurach stosowanych w organizacji projektowej. Sam model pozostałby „na papierze”, natomiast jego implementacja zostałaby uwzględniona jedynie jako zestaw reguł, procedur, integracji i ustawień w istniejących systemach informatycznych.

Niniejszy rozdział jest częścią szerszych badań dotyczących zastosowania podejścia obiektowego w zarządzaniu projektami, programami i portfelami projektów. W następnych krokach prezentowany powyżej model zostanie rozwinięty o obiekty będące częścią programu oraz zostanie zweryfikowane jego działanie w rzeczywistej organizacji biznesowej. Kolejnym krokiem powinna być budowa dedykowanego rozwiązania informatycznego wspierającego proponowane podejście.

---

<sup>27</sup> Przez określenie „papierowy” autor rozumie takie, które nie jest wspierane przez dedykowany system informatyczny, ale nadal może wykorzystywać e-mail czy cyfrowe przypomnienie ustawione w kalendarzu itp.

## Literatura

- Bergeron B. (2003), *Essentials of Shared Services*, John Wiley & Sons.
- Brady T., Davies A. (2014), *Managing Structural and Dynamic Complexity: A Tale of Two Projects*, „Project Management Journal”, Vol. 45(4), s. 21-38, doi: 10.1002/pmj.21434.
- Cartledge A., Hanna A., Rudd C., Macfarlane I., Windebank J., Rance S. (2007), *The IT Infrastructure Library. An Introductory Overview of ITIL® V3*, The UK Chapter of the itSMF.
- Cicmil S., Cooke-Davies T.J., Crawford L.H., Richardson K. (2009), *Exploring the Complexity of Projects: Implications of Complexity Theory for Project Management Practice*. Project Management Institute, Newtown Square, Pennsylvania 19073-3299 USA.
- Deloitte (2018), *The Symphonic Enterprise – Tech Trends 2018*, [https://www2.deloitte.com/content/dam/insights/us/articles/Tech-Trends-2018/4109\\_TechTrends-2018\\_FINAL.pdf](https://www2.deloitte.com/content/dam/insights/us/articles/Tech-Trends-2018/4109_TechTrends-2018_FINAL.pdf).
- Engwall M. (2003), *No Project Is an Island: Linking Projects to History and Context*, „Research Policy”, Vol. 32(5), s. 789-808, doi: 10.1016/S0048-7333(02)00088-4.
- Jaafari A., Manivong K. (1998), *Towards a Smart Project Management Information System*, „International Journal of Project Management”, Vol. 16(4), s. 249-265.
- Joslin R., Müller R. (2016), *The Relationship between Project Governance and Project Success*, „International Journal of Project Management”, Vol. 34(4), s. 613-626, doi: 10.1016/j.ijproman.2016.01.008.
- Maylor H., Turkulainen V. (2019), *The Concept of Organisational Projectification: Past, Present and Beyond?* „International Journal of Managing Projects in Business”, Vol. 12(3), s. 565-577, doi: 10.1108/IJMPB-09-2018-0202.
- Medina R., Medina A. (2015), *The Competence Loop: Competence Management in Knowledge-intensive, Project-intensive Organizations*, „International Journal of Managing Projects in Business”, Vol. 8(2), s. 279-299, doi: 10.1108/IJMPB-09-2014-0061.
- Müller R., Lecoivre L. (2014), *Operationalizing Governance Categories of Projects*, „International Journal of Project Management”, Vol. 32(8), doi: 10.1016/j.ijproman.2014.04.005.
- Riis E., Hellström M., Wikström K. (2019), *Governance of Projects: Generating Value by Linking Projects with Their Permanent Organisation*, „International Journal of Project Management”, Vol. 37(5), s. 652-667, doi: 10.1016/j.ijproman.2019.01.005.
- Rolstadås A., Schiefloe P.M. (2017), *Modelling Project Complexity*, „International Journal of Managing Projects in Business”, Vol. 10(2), s. 295-314.

- 
- Sydow J., Lindkvist L., Defillippi R. (2004), *Project-Based Organizations, Embeddedness and Repositories of Knowledge: Editorial*, „Organization Studies”, Vol. 25(9), s. 1475-1489.
- Tync Ł. (2012a), *Kontekst projektu w podejściu obiektywnym*, Prace Naukowe / Uniwersytet Ekonomiczny w Katowicach (Optymalizacja w praktyce), s. 117-134.
- Tync Ł. (2012b), *Podejście obiektywne i wzorzec Obserwator jako wsparcie dla zarządzania ryzykiem i komunikacją w projekcie*, Studia Ekonomiczne / Uniwersytet Ekonomiczny w Katowicach 97 (Modelowanie preferencji a ryzyko '12), s. 337-347.



## Część III

---

# UCZENIE MASZYNOWE



# Zastosowanie uczenia maszynowego w predykcji rezultatów meczów piłki nożnej

*Szymon Głowania<sup>1</sup>, Jan Kozak<sup>2</sup>*

## Wprowadzenie

Rosnąca dostępność narzędzi umożliwiających implementację algorytmów sztucznej inteligencji wpływa na coraz szersze jej zastosowanie oraz ciągle poszukiwanie możliwości optymalizacji. Znaczącym obszarem sztucznej inteligencji, który cieszy się obecnie dużą popularnością zarówno w świecie nauki, jak i biznesie, jest uczenie maszynowe (ang. machine learning).

Dużą popularnością cieszą się również różnego rodzaju rywalizacje sportowe, z którymi wiążą się ogromne fundusze i możliwości zysku. Coraz częściej zarówno firmy, jak i pasjonaci sięgają po metody sztucznej inteligencji w predykcji wyników rozgrywek sportowych lub zarządzaniu drużyną [Ćwikliński, Giełczyk, Choraś, 2021; Bejger, 2021]. W literaturze można znaleźć wiele publikacji dotyczących tego zagadnienia związanych z różnymi dyscyplinami i algorytmami. Brakuje jednak opracowania bazującego na kompleksowym podejściu do postawionego problemu z procesem odkrywania wiedzy z danych (ang. Knowledge Discovery in Databases, KDD).

Dokonując przeglądu literatury, można zaobserwować, iż znaczna większość publikacji i badań dotyczy angielskiej Premier League, która ma rozbudowane zbiory danych zawierające różnego rodzaju statystyki związane z tymi rozgrywkami [Baboota, Kaur, 2019; Eryarsoy, Delen, 2019; Joseph, Fenton, Neil, 2006; Schauburger, Groll, 2018; Schauburger, Groll, Tutz, 2016]. Pierwsze ligi innych krajów, jak również mniej popularne rozgrywki krajowe nie mają takiej liczby opracowań i pozostawiają nadal znaczne pole do badań.

---

<sup>1</sup> Uniwersytet Ekonomiczny w Katowicach, Wydział Informatyki i Komunikacji, Katedra Uczenia Maszynowego, e-mail: szymon.glowania@ue.katowice.pl.

<sup>2</sup> Uniwersytet Ekonomiczny w Katowicach, Wydział Informatyki i Komunikacji, Katedra Uczenia Maszynowego, e-mail: jan.kozak@ue.katowice.pl.

W dostępnych opracowaniach można również odnaleźć artykuły poruszające tematykę predykcji wyników w innych dyscyplinach sportu. Badania te dotyczą między innymi koszykówki [Pai, Chang Liao, Lin, 2017], rugby [McCabe, Trevathan, 2008], rzutu oszczepem [Maszczyk i in., 2014], futbolu amerykańskiego [Delen, Cogdell, Kasap, 2012] oraz tenisa [Wilkens, 2020].

Podczas doboru odpowiedniego algorytmu uczenia maszynowego często spotykanym podejściem jest porównywanie ze sobą (w różnych kombinacjach) wyników uzyskanych przez drzewa decyzyjne (ang. decision trees, DT), maszynę wektorów nośnych (ang. support vector machine, SVM) i lasy losowe (ang. random forest, RF), dlatego algorytmy te zostaną zastosowane w przygotowanym eksperymencie.

Głównym celem rozdziału jest wybranie algorytmów uczenia maszynowego pozwalających na poprawę predykcji wyników meczów piłkarskich w stosunku do przewidywań bukmacherów. Do weryfikacji, czy zaproponowany algorytm pozwala na poprawę wyników klasyfikacji, niezbędne jest opracowanie zbioru danych oraz wytrenowanie poszczególnych modeli. Postawiono hipotezę, iż zastosowanie opracowanego modelu opartego na prostej tabeli decyzyjnej pozwoli na poprawę predykcji wyników meczów piłkarskich polskiej PKO Ekstraklasy względem przewidywań bukmacherów. Ocena poprawy predykcji zostanie dokonana na podstawie uzyskanego zysku/ straty w symulacji zakładów bukmacherskich.

W kolejnej części rozdziału przedstawiono zagadnienia związane z dokonanym przeglądem literatury dotyczącej predykcji wyników sportowych. Następnie opisano zastosowane w przygotowanym rozwiązaniu algorytmy. Dalej zaprezentowano przeprowadzone eksperymenty i ich wyniki. W ostatniej sekcji krótko podsumowano uzyskane wyniki i przedstawiono dalsze wizje prac.

## **1.1. Uczenie maszynowe w predykcji wyników rozgrywek sportowych**

Duża popularność zawodów sportowych oraz związanego z nimi biznesu powoduje wzrost popularności metod analizy i predykcji wyników poszczególnych spotkań. Jedną z najczęściej analizowanych dyscyplin jest piłka nożna, w ramach której dużą popularnością cieszą się analizy angielskiej Premier League.

### 1.1.1. Predykcja wyników w piłce nożnej

Znacząca większość publikacji dotycząca predykcji wyników spotkań sportowych analizuje możliwości użycia uczenia maszynowego w predykcji wyników rozgrywek piłkarskich. Na czele wszystkich analiz znajduje się angielska Premier League, która jest analizowana z użyciem różnych metod uczenia maszynowego.

Jak zauważają w swojej pracy R. Bunker i F. Thabtah [2019], możliwości użycia sztucznych sieci neuronowych (ang. neural networks, NNs) do predykcji wyników meczów piłkarskich są znaczące. Przygotowana przez nich klasyfikacja odbywa się poprzez przyporządkowanie wyniku do jednej z trzech zaproponowanych klas: wygrana, remis, przegrana. Autorzy zaznaczają, iż ilość dostępnych danych mogących posłużyć jako zbiór uczący jest ogromna i może dotyczyć zarówno historycznych wyników drużyn, jak i informacji o zawodnikach czy konkretnym spotkaniu, a ich wybór jest jednym z kluczowych elementów. Praca zawiera również prezentację zaproponowanego podejścia do modelowania tego typu danych.

W artykule opracowanym przez K. Sujatha, T. Godhavari i N. Bhavanii [2018] została dokonana analiza dokładności przewidywania wyników meczów piłkarskich. Autorzy porównali opracowany model sztucznej sieci neuronowej z powszechnie używanym systemem FRES (ang. Football Result Expert System), który został stworzony dla poprawy dokonywanych przez ekspertów prognoz. Zaproponowane podejście pozwoliło na poprawę predykcji względem FRES, uwzględniało szereg zmiennych, a dokładność predykcji wzrastała wraz z aktualnością spotkań [Sujatha, Godhavari, Bhavanii, 2018].

W kolejnej pracy [Baboota, Kaur, 2019] poruszono temat przewidywania wyniku meczu, jak również liczby strzelonych bramek. R. Baboota i H. Kaur cytują wyniki badań z zastosowaniem innych podejść, między innymi:

- użycie sieci bayesowskiej (ang. Bayesian network) – A. Joseph, N. Fenton i M. Neil [2006],
- użycie próbkowania Monte Carlo łańcuchami Markowa w połączeniu z siecią bayesowską – H. Rue i O. Salvesen [2000],
- użycie logiki rozmytej (ang. fuzzy logic) – A.P. Rotshtein, M. Posner i A.B. Rakityanskaya [2005],
- użycie danych z mediów społecznościowych oraz tradycyjnych technik statystycznych – R.P. Schumaker, T.A. Jarmoszko i C.S. Labedz [2016].

Autorzy dokonują przekrojowego porównania możliwości predykcji dla piłki nożnej z użyciem różnych metod i zwracają uwagę, iż liga angielska cechuje się dużą zmiennością, co ma również duży wpływ na przewidywanie wyników. Uzyskany model cechował się znaczną dokładnością, jednak w porównaniu z zastosowanymi przez bukmacherów rozwiązaniami nie pozwalał na uzyskanie lepszych wyników.

E. Eryarsoy i D. Delen [2019] dokonują analizy na podstawie Turkish Super League (wyniki z lat 2007-2017) dotyczącej predykcji wyników meczów piłkarskich. Predykcja jest dokonywana na dwóch płaszczyznach wygrana/remis/ przegrana oraz drużyna zdobędzie punkty/ drużyna nie zdobędzie punktów. Szczególną uwagę poświęcono algorytmom: lasy losowe, drzewa decyzyjne, naiwny klasyfikator bayesowski (ang. naive Bayes classifier), zespół klasyfikatorów, maszyna wektorów nośnych, metoda  $k$  najbliższych sąsiadów (ang.  $k$  nearest neighbours,  $k$ -nn) oraz sieci neuronowe. Najwyższą dokładność na badanym zbiorze uzyskały lasy losowe.

W pracy Razali i in. [2017] autorzy podejmują się klasyfikacji wyników meczów piłkarskich angielskiej Premier League z użyciem sieci bayesowskiej (ang. Bayesian Networks), a uzyskane przez nich wyniki są porównywalne z tymi uzyskanymi w pracy Joseph, Fenton, Neil [2006]. W analizie zostały użyte dane dotyczące 3 sezonów rozgrywek: 2010-2011, 2011-2012 i 2012-2013. Otrzymane wyniki cechowały się większą dokładnością niż pierwotnie uzyskane rozwiązanie.

G. Schauburger, A. Groll i G. Tutz [2016] dokonują analizy możliwości przewidywania wyniku meczu z zastosowaniem trzech odmian modelu Bradleya-Terry'ego, który odpowiednio rozbudowują dla poprawy wyniku. Oprócz zastosowania statystycznego podejścia do predykcji celem autorów jest identyfikacja głównych zmiennych wpływających na sukces drużyny. Dokonują oni analizy zbieranych i możliwych do użycia zmiennych dotyczących poszczególnych statystyk zespołów związanych np. z procentem posiadania piłki lub efektem rozgrywek na własnym stadionie. Jest to jeden z niewielu artykułów, w którym autorzy prowadzą badania na danych dotyczących meczów piłkarskich (sezon 2015/16) niemieckiej Bundesligi. Uzyskana dokładność przewyższała prognozy bukmacherów.

### 1.1.2. Uczenie maszynowe w predykcji wyników innych rozgrywek sportowych

Mimo znaczącej przewagi publikacji dotyczących analizy możliwości zastosowania uczenia maszynowego w predykcji wyników piłki nożnej można również znaleźć opracowania dotyczące innych sportów. Rozgrywkami cieszącymi się znaczną popularnością są między innymi: koszykówka, rzut oszczepem, futbol amerykański i wiele innych.

Podczas doboru odpowiedniego algorytmu uczenia maszynowego najczęściej spotykanym podejściem jest porównywanie ze sobą wyników uzyskanych przez drzewa decyzyjne, maszynę wektorów nośnych i sztuczne sieci neuronowe. W większości badań to ostatnie z wymienionych podejść uzyskuje największą dokładność, jednak jak można zaobserwować na analizie dotyczącej możliwości predykcji wyniku (wygrana/przegrana) dla meczów futbolu amerykańskiego zaprezentowanej w pracy Delen, Cogdell, Kasap [2012], sytuacja ta nie jest regułą. Dla analizowanych danych drzewa decyzyjne osiągnęły większą dokładność niż sieci neuronowe czy maszyna wektorów nośnych.

W opracowaniu Maszczyk i in. [2014] autorzy dokonali porównania możliwości predykcji długości rzutu oszczepem z zastosowaniem statystycznego modelu regresji liniowej (ang. linear regression) oraz sztucznej sieci neuronowej. Najlepsza z przygotowanych sieci wykazała dla 20 prób sumę błędów na poziomie 16,7 metra, przy sumie błędów 29,45 metra dla modelu regresji. W analizie użyto wyniki 116 oszczepników w wieku około 18 lat. Najlepsze rozwiązanie w postaci sztucznej sieci neuronowej miało strukturę (4–3–2–1), a cechy pozwalające na uzyskanie takiego wyniku to między innymi: postawa podczas rzutu, siła ramion, siła mięśni brzucha oraz siła chwytu.

Jak zaprezentowano w pracy Leung i Joseph [2014], podczas predykcji college football dzięki zastosowaniu przez autorów nowego podejścia osiągnięto przewagę dokładności predykcji nad standardowymi rozwiązaniami związanymi z sieciami neuronowymi, maszyną wektorów nośnych i drzewami decyzyjnymi. Zaproponowane podejście opiera się na 4-wymiarowej przestrzeni cech poszczególnych drużyn, na podstawie których jest tworzony wskaźnik zastosowany do predykcji wyniku.

W kolejnej pracy [Kahn, 2003] przedstawiono zastosowanie sztucznych sieci neuronowych w przewidywaniu wyników meczów NFL (ang. National Football League). J. Kahn użył wytrenowany na 192 spotkaniach model i uzyskał wyższą dokładność predykcji niż eksperci z ESPN.com. Wyniki zostały przetestowane na 16 spotkaniach z 14. i 15. tygodnia rozgrywek.

Cai i in. [2019] prezentują własne podejście do predykcji wyników meczów koszykówki z zastosowaniem uczenia maszynowego. W ramach przeprowadzonych eksperymentów wykazują przewagę dokładności predykcji algorytmu AdaBoost między innymi nad naiwnym klasyfikatorem bayesowskim, modelem Markova, regresją logistyczną (ang. logistic regression), maszyną wektorów nośnych czy sztucznymi sieciami neuronowymi.

## **1.2. Algorytmy uczenia maszynowego**

Klasyfikacja jako jedna z metod pozwalająca na eksplorację danych polega na stworzeniu modelu, dzięki któremu możliwe będzie przyporządkowanie nowego obiektu do jednej z predefiniowanych klas [Kozak, 2019]. Wśród najpopularniejszych metod klasyfikacji można wyróżnić klasyfikację bayesowską, metodę  $k$ -najbliższych sąsiadów, drzewa decyzyjne, lasy losowe, maszynę wektorów nośnych czy sztuczne sieci neuronowe.

### **1.2.1. Drzewa decyzyjne**

Drzewa decyzyjne (ang. decision trees, DT) są algorytmami służącymi do klasyfikacji, które swoją strukturą przypominają drzewo, a ich główne elementy to gałęzie i węzły. Drzewa decyzyjne są skierowanymi grafami acyklicznymi, węzły są tworzone z wierzchołków, a gałęzie to powiązania między nimi. Każdy z liści jest etykietą odpowiedniej kategorii, natomiast węzeł jest testem atrybutów warunkowych, zaś krawędź to wynik testu. Korzeniem drzewa jest wierzchołek bez rodzica, a wierzchołki bez potomków to liście.

Klasyfikacja atrybutu z zastosowaniem drzewa polega na jego przesuwaniu od głównego węzła (korzenia) do liści poprzez węzły, dla których spełnia on testy. Proces ten zostaje zakończony, gdy atrybut dotrze do liścia. Po dotarciu do liścia atrybut otrzymuje etykietę klasy, która jest do niego przypisana. W pierwszej kolejności zostaje wybrany atrybut reprezentujący korzeń drzewa. Następnie jest tworzona gałąź dla każdej możliwej wartości atrybutu. Dalej następuje podział zestawu danych na podzbiory względem wartości atrybutu. W kolejnych krokach rekurencyjnie są wybierane cechy dla kanału z użyciem węzła odpowiedniego podzbioru atrybutów o zbliżonych wartościach. W literaturze można znaleźć wiele algorytmów pozwalających tworzyć drzewa decyzyjne, między innymi ID3, C4.5, C5.0 lub CART [Kozak, 2019].

Algorytm CART (ang. Classification and Regression Trees, CART) został opracowany przez L. Breimana i in. w 1984 roku. Dla tego algorytmu zostały zaproponowane dwa kryteria podziału: Giniego oraz podziału na dwie części. Z użyciem kryterium podziału jest szukany test, który najlepiej podzieli dane w węźle na dwie maksymalnie jednorodne części. Kryterium Giniego opiera się na indeksie Giniego, który jest miarą koncentracji zmiennej losowej. Jego działanie zmierza do podziału zbioru na jednorodne podzbiory w węzłach potomnych. Kryterium podziału na dwie części (ang. twoing rule) w głównej mierze opiera swoje działanie na podziale danych na dwa maksymalnie zbliżone do siebie liczebnie zbiory. W przypadku tego podejścia jednorodność zbiorów schodzi na dalszy plan [Breiman i in., 1984].

### **1.2.2. Maszyna wektorów nośnych**

Maszyna wektorów nośnych (ang. support vector machine, SVM) jest techniką uczenia maszynowego, której celem jest klasyfikacja nowego elementu do jednej z dwóch predefiniowanych klas. Do trenowania jest używany zbiór uczący z oznaczonymi przypadkami. Stworzony model zawiera informację o granicy rozdzielającej przypadki do poszczególnych klas.

Głównym zadaniem algorytmu jest odnalezienie hiperpłaszczyzny, która w optymalny sposób rozdzieli przypadki do poszczególnych klas. Dążenie do optimum uzyskuje się poprzez otrzymanie możliwie najszerzego marginesu pomiędzy klasami, co jest związane ze zwiększaniem odległości punktów od hiperpłaszczyzny. Pierwsza prezentacja SVM została dokonana przez V. Vapnika [Cortes, Vapnik, 1995]. SVM posiada trzy główne typy klasyfikatorów: liniowy, nieliowy i C-SVC.

### **1.2.3. Zespoły klasyfikatorów**

Zespoły klasyfikatorów (ang. Ensemble Method, EM) są nazywane rodziną klasyfikatorów. Podejście to polega na połączeniu ze sobą pojedynczych klasyfikatorów, które poprzez wielokrotne uruchomienie algorytmu budują wiele hipotez na podstawie zróżnicowanych próbek danych. Decyzja dotycząca końcowej klasyfikacji jest podejmowana z uwzględnieniem wszystkich utworzonych klasyfikatorów. Prowadzone badania teoretyczne oraz empiryczne potwierdzają przewagę dokładności predykcji zespołów klasyfikatorów nad zastosowaniem pojedynczych klasyfikatorów [Hansen, Salamon, 1990; D'zeroski, Zenko, 2004; Opitz, Maclin, 1999].

W literaturze dominują dwa podejścia dotyczące EM. Pierwsze z nich opiera się na generowaniu zestawu hipotez niezależnie, dzięki czemu uzyskany zbiór jest dokładny i zróżnicowany. Kolejne podejście polega na tworzeniu hipotez zależnych, aby każda z nich dzięki odpowiedniej wadze pozwalała na dobre dopasowanie do danych. Rodzina klasyfikatorów podejmuje decyzję za pomocą głosowania prostego. Pojedyncze klasyfikatory przydzielają kategorię dla próbki, a następnie zostaje wybrana klasa, która uzyskała najwięcej głosów. Najpopularniejsze rodziny klasyfikatorów to boosting, lasy losowe i bagging [Kozak, 2019].

### 1.2.3.1. Boosting i AdaBoost

Boosting to dyskretny adaptacyjny algorytm uczący zaproponowany przez R.E. Schapire’a. Jego idea opiera się na stworzeniu „mocnego i złożonego klasyfikatora” z użyciem „słabych i prostych klasyfikatorów”. Pierwotne podejście zostało udoskonalone przez Y. Freunda i R.E. Schapire’a i opublikowane w pracach Freund i Schapire [1996; 1997] jako algorytm AdaBoost (ang. Adaptive Boosting, AB).

Algorytm, rozpoczynając działanie, przydziela elementom zbioru uczącego wagi ( $1/n$ ) stanowiące prawdopodobieństwo wylosowania elementu do pseudo-próbki ( $n$  – liczba elementów zbioru). Dokonywana jest klasyfikacja i ocena poprawności przypisania elementów do klas. Następnie wagi zostają zaktualizowane. Elementy pierwotnie błędnie zaklasyfikowane otrzymują wyższe wagi, co pozwala na zwiększenie prawdopodobieństwa wylosowania elementu. Zwiększenie jest uzależnione od wielkości błędu klasyfikacji, czyli sumy wag elementów, które były błędnie sklasyfikowane [Kozak, 2019].

### 1.2.3.2. Lasy losowe

Lasy losowe (ang. random forest, RF) to zaproponowana przez L. Breimana metoda konstruowania zespołu klasyfikatorów. Lasy losowe składają się z wielu drzew decyzyjnych, które na drodze głosowania prostego podejmują decyzję dotyczącą końcowej klasyfikacji.

W pierwszym kroku algorytmu są generowane pseudoprobki, które zostaną zastosowane do tworzenia drzew decyzyjnych. Losowanych jest  $n$  elementów (ze zwracaniem) do pseudoprob ze zbioru  $n$ -elementowego. Każdy z elementów ma jednakowe prawdopodobieństwo wylosowania  $1/n$ . Z użyciem utworzonych pseudoprob są generowane drzewa decyzyjne. Proces zostaje zakończony

w chwili, gdy dla każdego z drzew w węźle pozostaną elementy należące do jednej klasy. Podczas tworzenia drzewa każda reguła dzielenia jest stosowana niezależnie dla różnych podzbiorów atrybutów [Breiman, 2001].

## **1.3. Eksperyment**

Zaprezentowane w rozdziale rozwiązanie jest kompleksowym podejściem KDD (ang. Knowledge Discovery in Databases). Niezbędne dane zostały zebrane, przetworzone, a następnie zapisane w postaci tabeli decyzyjnej. W kolejnym kroku przeprowadzono analizę dostępnych algorytmów uczenia maszynowego. Porównano dostarczane przez algorytmy wyniki pod względem dokładności klasyfikacji i możliwości uzyskania przewagi nad predykcją opracowaną przez bukmachera.

Poniżej przedstawiono obserwacje powstałe w wyniku przeprowadzonego eksperymentu, którego celem była weryfikacja zaproponowanego podejścia. Wnioski te pozwoliły na potwierdzenie hipotezy badawczej o możliwości skutecznego zastosowania prostej tabeli decyzyjnej do predykcji wyników spotkań piłkarskich i wyborze algorytmu pozwalającego na uzyskanie największej przewagi nad predykcjami bukmachera.

### **1.3.1. Zbiór danych i tworzenie eksperymentu**

Pierwszą fazą opracowania rozwiązania jest przygotowanie danych. Etap ten jest zależny od wybranej ligi i dostępnych źródeł. Dane przechodzą proces selekcji i transformacji. W konsekwencji otrzymane dane mogą zostać zastosowane do uczenia klasyfikatorów. Dane ograniczono do 5 atrybutów warunkowych i decyzji (rezultatów spotkań).

Zaproponowane podejście oparto na łatwo dostępnych danych. Ze względu na dobór popularnych i ogólnodostępnych informacji rozwiązanie to może być w prosty sposób skalowane na inne ligi piłkarskie, w tym zagraniczne. Zestaw danych ograniczono do pięciu atrybutów:

- (1) aktualna pozycja drużyny 1 – w tabeli ligowej,
- (2) aktualna pozycja drużyny 2,
- (3) liczba punktów drużyny 1,
- (4) liczba punktów drużyny 2,
- (5) różnica pomiędzy punktami drużyny 1 i 2.

Wyznaczono także trzy klasy decyzyjne, które reprezentują rezultat spotkania: 1 – zwycięstwo gospodarza; 0 – remis; 2 – zwycięstwo drużyny przyjezdnej.

Na podstawie wniosków wynikających z przeglądu literatury i wstępnych eksperymentów zdecydowano o tworzeniu modeli uczenia maszynowego w oparciu o wyniki wszystkich spotkań rozegranych od 6. kolejki. Uwzględnienie 5 pierwszych kolejek skutkowało obniżeniem dokładności predykcji dla modelu. Do predykcji zostały zastosowane kolejki od 16. do 30. sezonu 2019/2020 polskiej PKO Ekstraklasy. Mecze te były rozgrywane pomiędzy 22 listopada 2019 roku a 14 czerwca 2020 roku i obejmowały 120 zdarzeń.

Doświadczenia były prowadzone w dwóch etapach: trenowania i testowania modeli. Zebrane dane zostały podzielone na dwa zbiory: treningowy i testowy. W ramach zbioru treningowego znalazły się wyniki wszystkich spotkań od 6. kolejki sezonów od 2013/2014 do 2018/2019 (6 sezonów). W ramach zbioru testowego dostępne były wszystkie rozgrywki w kolejkach od 16. do 30. sezonu 2019/2020. Zastosowane podejście pozwoliło na ocenę jakości klasyfikacji oraz odwzorowanie rzeczywistej sytuacji, gdy do uczenia służą dane historyczne, a klasyfikacja odbywa się na bieżących danych. Zaproponowane podejście umożliwiło przeprowadzenie symulacji pozwalającej na porównanie z bukmacherem. Do oceny wyników użyto miary dokładności klasyfikacji i porównania do bukmachera z zastosowaniem wirtualnego portfela.

Utworzony zbiór treningowy zawiera 1536 przypadków (6 sezonów po 32 kolejki, a w nich po 8 spotkań), natomiast zbiór testowy zawiera 120 przypadków (15 kolejek po 8 spotkań). Podział na klasy decyzyjne został zaprezentowany w tabeli 1 i wskazuje na podobny rozkład klas decyzyjnych w każdym ze zbiorów danych.

**Tabela 1.** Udział przypadków dla klas decyzyjnych w zbiorze treningowym i testowym

| Klasy decyzyjne | Liczba przypadków<br>w zbiorze treningowym | Liczba przypadków<br>w zbiorze testowym |
|-----------------|--------------------------------------------|-----------------------------------------|
| 1               | 689 (45%)                                  | 61 (51%)                                |
| 2               | 431 (28%)                                  | 31 (26%)                                |
| 0               | 416 (27%)                                  | 28 (23%)                                |

Źródło: Opracowanie własne.

Jako miarę oceny jakości modeli użyto dokładności klasyfikacji (ang. accuracy). Miara ta została obliczona według wzoru 1. Do obliczenia miary zastosowano macierz błędów, która została zaprezentowana w tabeli 2.

$$\text{dokładność klasyfikacji} = \frac{n_{00} + n_{11} + n_{22}}{n_{00} + n_{01} + n_{02} + n_{10} + n_{11} + n_{12} + n_{20} + n_{21} + n_{22}} \quad (1)$$

$n_{00}$  – obserwacje z klasy 0 zaklasyfikowane poprawnie przez model do klasy 0,  
 $n_{11}$  – obserwacje z klasy 1 zaklasyfikowane poprawnie przez model do klasy 1,  
 $n_{22}$  – obserwacje z klasy 2 zaklasyfikowane poprawnie przez model do klasy 2,  
 $n_{01}$  – obserwacje z klasy 0 zaklasyfikowane niepoprawnie przez model do klasy 1,  
 $n_{02}$  – obserwacje z klasy 0 zaklasyfikowane poprawnie przez model do klasy 2,  
 $n_{10}$  – obserwacje z klasy 1 zaklasyfikowane niepoprawnie przez model do klasy 0,  
 $n_{12}$  – obserwacje z klasy 1 zaklasyfikowane poprawnie przez model do klasy 2,  
 $n_{20}$  – obserwacje z klasy 2 zaklasyfikowane niepoprawnie przez model do klasy 0,  
 $n_{21}$  – obserwacje z klasy 2 zaklasyfikowane poprawnie przez model do klasy 1.

**Tabela 2.** Macierz błędów

|              |         | Obserwacja |          |          |
|--------------|---------|------------|----------|----------|
|              |         | klasa 0    | klasa 1  | klasa 2  |
| Klasyfikacja | klasa 0 | $n_{00}$   | $n_{01}$ | $n_{02}$ |
|              | klasa 1 | $n_{10}$   | $n_{11}$ | $n_{12}$ |
|              | klasa 2 | $n_{20}$   | $n_{21}$ | $n_{22}$ |

Źródło: Opracowanie własne.

W celu porównania wybranych metod przeprowadzono eksperymenty dla poszczególnych klasyfikatorów. Na podstawie przeglądu literatury do eksperymentu wybrano następujące algorytmy: drzewa decyzyjne (CART), maszynę wektorów nośnych, lasy losowe i algorytm AdaBoost. Dokładne wyniki eksperymentów zaprezentowano w kolejnej sekcji.

### 1.3.2. Wyniki

Wyniki oceny jakości klasyfikacji dla wybranych algorytmów zaprezentowano w tabeli 3. Jako miara oceny jakości modeli została wybrana opisana powyżej dokładność klasyfikacji (wzór 1).

**Tabela 3.** Dokładność klasyfikacji poszczególnych algorytmów

| Podejście                                            | Dokładność klasyfikacji |
|------------------------------------------------------|-------------------------|
| Losowe                                               | 33,33%                  |
| Drzewa decyzyjne CART (bez przycinania), DT0         | 37,50%                  |
| Drzewa decyzyjne CART (przcięte do 3. poziomu), DT   | 51,67%                  |
| Maszyna wektorów nośnych, SVM                        | 50,83%                  |
| Lasy losowe (100 drzew przciętych do 3. poziomu), RF | 54,17%                  |
| AdaBoost, AB                                         | 52,50%                  |

Źródło: Opracowanie własne.

Wszystkie z zaproponowanych podejść pozwoliły na uzyskanie lepszych wyników od podejścia losowego. Największą dokładność uzyskało podejście oparte na lasach losowych. Dokładność klasyfikacji dla tego algorytmu na zbiorze testowym wyniosła 54,17%. Algorytmem, który uzyskał również zadowalający wynik (52,50%), jest AdaBoost.

Następnie dla uzyskanych predykcji przeprowadzono symulację polegającą na zagraniu na każde z klasyfikowanych zdarzeń (symulacja została przeprowadzona dla każdego z testowanych algorytmów). Dla ułatwienia obliczeń każde zdarzenie zostało zagrane za kwotę 10 PLN. W związku z opodatkowaniem gier hazardowych (podatek 12% od stawki na zakładzie) kurs dla danego spotkania był mnożony przez 8,8. Zastosowano rzeczywiste kursy z dnia wykonania klasyfikacji.

**Tabela 4.** Symulacja zakładów bukmachera

| Podejście | Suma wkładów (PLN) | Suma wygranych (PLN) | Bilans (PLN) | Procentowy zysk/strata |
|-----------|--------------------|----------------------|--------------|------------------------|
| DT0       | 1200,00            | 1148,17              | -51,83       | -4,32%                 |
| DT        | 1200,00            | 1246,17              | +46,17       | 3,85%                  |
| SVM       | 1200,00            | 1214,94              | +14,94       | 1,25%                  |
| RF        | 1200,00            | 1268,53              | +68,53       | 5,71%                  |
| AB        | 1200,00            | 1250,41              | +50,41       | 4,20%                  |

Źródło: Opracowanie własne.

W tabeli 4 zaprezentowano wyniki przeprowadzonej symulacji. Algorytmy DT, SVM, RF i AB uzyskały dodatni bilans sumy wkładów brutto do wygranych netto. Jedynym algorytmem, który uzyskał ujemny bilans, jest DT0 – drzewa decyzyjne CARD (bez przycinania). W przeprowadzonej symulacji, podobnie jak w przypadku dokładności klasyfikacji, najlepszy wynik uzyskały lasy losowe RF. Przy wkładzie 1200 PLN pozwoliłyby one na uzyskanie wygranych o sumie bliskiej 1269 PLN, co stanowi procentowy przyrost na poziomie 5,71%.

Z przeprowadzonych eksperymentów i symulacji wynika, iż jest możliwe zastosowanie algorytmów uczenia maszynowego w celu predykcji wyników meczów piłkarskich. Jak zaobserwowano na przykładzie wyników uzyskanych na polskiej PKO Ekstraklasie, przy stosunkowo prostej tabeli decyzyjnej możliwe jest otrzymanie satysfakcjonujących wyników. Algorytmami, które uzyskały najlepsze wyniki w ramach badanych przypadków, są lasy losowe i algorytm AdaBoost.

## **Podsumowanie i dalsze prace**

Głównym celem rozdziału było wybranie algorytmów uczenia maszynowego pozwalających na poprawę predykcji wyników meczów piłkarskich w stosunku do przewidywań bukmacherów. Realizacja celu pozwoliła na potwierdzenie hipotezy badawczej o możliwości skutecznego zastosowania prostej tabeli decyzyjnej do predykcji wyników spotkań piłkarskich i wyborze algorytmu pozwalającego na uzyskanie największej przewagi nad predykcjami bukmachera. Założone cele zostały zrealizowane, a zaproponowane podejście pozwala na jego skalowanie m.in. na ligi zagraniczne. Uzyskane wyniki przewyższały predykcję losową. Najlepiej pod względem dokładności klasyfikacji i zysku w przeprowadzonej symulacji na zakładach bukmacherskich wypadły lasy losowe, a następny w kolejności był algorytm AdaBoost.

W przypadku badanego zbioru zaobserwowano bardzo małą predykcję dla klasy 0 (remis). Jest to znaczącym ograniczeniem dla badanych algorytmów. Zarówno w danych treningowych, jak i testowych procentowy udział przypadków należących do tej klasy znajdował się na poziomie klasy 2 (wygrana drużyny przyjezdnej), jednak predykcja tego typu przypadków była znikoma. Jest to pole do dalszych prac i udoskonalania rozwiązania.

Możliwe są również dalsze prace związane z rozbudową zbioru danych o nowe atrybuty, a także uwzględnienie nowych mechanizmów i wiedzy dziedzinowej. W ramach kontynuacji podejścia możliwe jest zastosowanie sztucznych sieci neuronowych oraz heterogenicznych zespołów klasyfikatorów w predykcji wyników spotkań. Rozwiązania te mogą pozwolić na dalszą poprawę predykcji.

## Literatura

- Baboota R., Kaur H. (2019), *Predictive Analysis and Modelling Football Results Using Machine Learning Approach for English Premier League*, "International Journal of Forecasting", Vol. 35(2), s. 741-755.
- Bejger S. (2021), *Optimization of the Structure of the Speedway Team in the Presence of Regulation and Financial Constraints – Application of Machine Learning and Integer Programming to the Polish League Case*, "Journal of Physical Education and Sport", doi:10.7752/jpes.2021.s2123, s. 990-998.
- Breiman L. (2001), *Random Forests*, "Machine Learning", Vol. 45(1), s. 5-32.
- Breiman L., Friedman J., Stone C., Olshen R. (1984), *Classification and Regression Trees*, Chapman & Hall, New York.
- Bunker R.P., Thabtah F. (2019), *A Machine Learning Framework for Sport Result Prediction*, "Applied Computing and Informatics", Vol. 15(1), s. 27-33.
- Cai W., Yu D., Wu Z., Du X., Zhou T. (2019), *A Hybrid Ensemble Learning Framework for Basketball Outcomes Prediction*, "Physica A: Statistical Mechanics and Its Applications", Vol. 528, s. 1-8.
- Cortes C., Vapnik V. (1995), *Support-Vector Networks*, "Machine Learning", Vol. 20(3), s. 273-297.
- Ćwiklinski B., Gielczyk A., Choraś M. (2021), *Who Will Score? A Machine Learning Approach to Supporting Football Team Building and Transfers*, "Entropy", Vol. 23(1), s. 1-12.
- Dżeroski S., Zenko B. (2004), *Is Combining Classifiers with Stacking Better Than Selecting the Best One?* "Machine Learning", Vol. 54(3), s. 255-273.
- Delen D., Cogdell D., Kasap N. (2012), *A Comparative Analysis of Data Mining Methods in Predicting NCAA Bowl Outcomes*, "International Journal of Forecasting", Vol. 28(2), s. 543-552.
- Eryarsoy E., Delen D. (2019), *Predicting the Outcome of a Football Game: A Comparative Analysis of Single and Ensemble Analytics Methods*, Proceedings of the 52nd Hawaii International Conference on System Sciences.
- Freund Y., Schapire R.E. (1996), *Experiments with a New Boosting Algorithm*, "ICML", Vol. 96, s. 148-156.
- Freund Y., Schapire R.E. (1997), *A Decision-theoretic Generalization of On-line Learning and an Application to Boosting*, "Journal of Computer and System Sciences", Vol. 55(1), s. 119-139.
- Hansen L.K., Salamon P. (1990), *Neural Network Ensembles*, "IEEE Transactions on Pattern Analysis and Machine Intelligence", Vol. 12(10), s. 993-1001.
- Joseph A., Fenton N.E., Neil M. (2006), *Predicting Football Results Using Bayesian Nets and Other Machine Learning Techniques*, "Knowledge-Based Systems", Vol. 19(7), s. 544-553.

- Kahn J. (2003), *Neural Network Prediction of NFL Football Games*, World Wide Web electronic publication, s. 9-15.
- Kozak J. (2019), *Decision Tree and Ensemble Learning Based on Ant Colony Optimization*, Springer.
- Leung C.K., Joseph K.W. (2014), *Sports Data Mining: Predicting Results for the College Football Games*, "Procedia Computer Science", Vol. 35, s. 710-719.
- Maszczyk A., Gołaś A., Pietraszewski P., Rocznio R., Zajac A., Stanula A. (2014), *Application of Neural and Regression Models in Sports Results Prediction*, "Procedia-Social and Behavioral Sciences", Vol. 117, s. 482-487.
- McCabe A., Trevathan J. (2008), *Artificial Intelligence in Sports Prediction*, Fifth International Conference on Information Technology: New Generations (ITNG 2008), s. 1194-1197.
- Opitz D., Maclin R. (1999), *Popular Ensemble Methods: An Empirical Study*, "Journal of Artificial Intelligence Research", Vol. 11, s. 169-198.
- Pai P.-F., Chang Liao L.-H., Lin K.-P. (2017), *Analyzing Basketball Games by a Support Vector Machines with Decision Tree Model*, "Neural Computing and Applications", Vol. 28(12), s. 4159-4167.
- Razali N., Mustapha A., Yatim F.A., Ab Aziz R. (2017), *Predicting Football Matches Results Using Bayesian Networks for English Premier League (EPL)*, IOP Conference series: Materials Science and Engineering, Vol. 226, s. 1-6, IOP Publishing.
- Rotshtein A.P., Posner M., Rakityanskaya A. (2005), *Football Predictions Based on a Fuzzy Model with Genetic and Neural Tuning*, "Cybernetics and Systems Analysis", Vol. 41(4), s. 619-630.
- Rue H., Salvesen O. (2000), *Prediction and Retrospective Analysis of Soccer Matches in a League*, "Journal of the Royal Statistical Society: Series D (The Statistician)", Vol. 49(3), s. 399-418.
- Schauberger G., Groll A. (2018), *Predicting Matches in International Football Tournaments with Random Forests*, "Statistical Modelling", Vol. 18(5-6), s. 460-482.
- Schauberger G., Groll A., Tutz G. (2016), *Modeling Football Results in the German Bundesliga Using Match-specific Covariates*, "Ludwig Maximilians Universität", Technical Reports, No. 197, doi: 10.5282/ubm/epub.29390, s. 1-21.
- Schumaker R.P., Jarmoszko A.T., Labeledz Jr C.S. (2016), *Predicting Wins and Spread in the Premier League Using a Sentiment Analysis of Twitter*, "Decision Support Systems", Vol. 88, s. 76-84.
- Sujatha K., Godhavari T., Bhavani N.P. (2018), *Football Match Statistics Prediction Using Artificial Neural Networks*, "International Journal of Mathematical and Computational Methods", Vol. 3, s. 1-8.
- Wilkens S. (2020), *Sports Prediction and Betting Models in the Machine Learning Age: The Case of Tennis*, "Journal of Sports Analytics", (Preprint), s. 1-19.



## Rozwój wyjaśniania modeli uczenia maszynowego a zaufanie do sztucznej inteligencji

*Szymon Głowania<sup>1</sup>, Bogna Zacny<sup>2</sup>, Grzegorz Dzikowski<sup>3</sup>*

### 2.1. Zaufanie i nieufność – ewolucja terminu

Przedmiotem niniejszego rozdziału jest wyjaśnienie pojęcia zaufania i nieufności w odniesieniu do algorytmów uczenia maszynowego. Przedstawiono w nim ewolucję terminu „zaufanie” i „nieufność” w naukach społecznych.

Na początku przyjrzymy się znaczeniu słowa „zaufanie”. W pierwszych publikacjach opisujących ten termin zaufanie było określane jako zaufanie jednostki do intencji i motywów innej osoby oraz opierało się na szczerości wypowiedzi [Mellinger, 1956]. Mając nadzieję na szczerą wymianę informacji, jesteśmy skłonni powierzyć osobiste informacje, narażając się w ten sposób na naruszenie naszego zaufania [Read, 1962]. Zaufanie zdefiniowano również jako zachowanie zgodne z naszymi oczekiwaniami. Osoba, której ufamy lub której zamierzamy zaufać, powinna się zachowywać w poszczególnych sytuacjach zgodnie z naszymi oczekiwaniami. W latach 90. znaczenie słowa „zaufanie” wyewoluowało z definicji opartej na motywach oraz intencjach do definicji opartej na zachowaniu. Zostało zdefiniowane jako gotowość do bycia podatnym na działanie innej osoby. Gotowość ta opiera się na oczekiwaniu, że wykonane zostanie działanie, niezależnie od tego, czy efekty będą monitorowane przez ufającego [Mayer, Davis, Schoorman, 1995]. Nieufność badacze zdefiniowali jako oczekiwanie, że osoby trzecie nie będą działać w naszym najlepszym interesie, a wręcz mogą działać szkodliwie [Barber, 1983]. W późnych latach 90.

---

<sup>1</sup> Uniwersytet Ekonomiczny w Katowicach, Wydział Informatyki i Komunikacji, Katedra Ucznia Maszynowego, e-mail: szymon.glowania@ue.katowice.pl.

<sup>2</sup> Uniwersytet Ekonomiczny w Katowicach, Wydział Informatyki i Komunikacji, Katedra Inżynierii Wiedzy, e-mail: bogna.zacny@ue.katowice.pl.

<sup>3</sup> Uniwersytet Ekonomiczny w Katowicach, Wydział Informatyki i Komunikacji, Katedra Ucznia Maszynowego, e-mail: grzegorz.dzikowski@ue.katowice.pl.

badacze zasugerowali, że zaufanie nie jest przeciwieństwem nieufności [Lewicki, McAllister, Bies, 1998]. Istnieją odrębne czynniki przyczyniające się do wzrostu i spadku zaufania czy nieufności. Badania wykazały, że zaufanie nie powinno być odbierane jako „coś dobrego”, a nieufność jako „coś złego”, ponieważ są to odmienne kategorie, można komuś ufać i nie ufać jednocześnie.

Wysokie zaufanie koreluje z takimi czynnikami, jak nadzieja, wiara, pewność, bezpieczeństwo i inicjatywa. Na wysoki poziom nieufności oddziałują wysokie wartości czynników, takich jak strach, sceptycyzm, cynizm, wojowniczość oraz czujność. Przedstawiając na osi współrzędnych dwa wymiary: ufność i nieufność, otrzymujemy 4 sytuacje relacji społecznych wyznaczone przez ćwiartki układu współrzędnych. Każda z ćwiartek przedstawia odmienną relację społeczną i zażyłość między osobami. Pierwsza relacja o niskim poziomie zaufania i nieufności przedstawia sytuację, gdzie osoby nie mają powodów do odczuwania pewności siebie, ale też nie czują się zagrożone. Relacje są sporadyczne, a ich istnienie nie narusza prywatności którejkolwiek ze stron lub nie sugeruje istnienia żadnej bliskości czy intymności. Druga relacja wysokiego zaufania i niskiej nieufności przedstawia sytuację, gdy jedna osoba dysponuje powodami, by mieć zaufanie do drugiej osoby i nie ma powodu, by podejrzewać ją o nieczyste intencje. Osoby te realizują wspólne cele. Poszukują nowych sposobów na stały rozwój stworzonej relacji. Strony tego typu relacji będą często wchodzić w interakcje i pielęgnować istniejące zaufanie. Trzecia relacja przedstawia warunki wysokiej nieufności i niskiego zaufania. Jedna osoba nie ma powodu, by ufać drugiej osobie i ma wystarczający powód do zachowania czujności i ostrożności. Warunki takie uniemożliwiają poprawną relację, gdyż odzwierciedlają dużą liczbę negatywnych doświadczeń. Ostatnia relacja przedstawia sytuację wysokiego zaufania i dużej nieufności. Opisywana jest w niej sytuacja, gdy jedna osoba w pewnych okolicznościach ma powody, aby mieć zaufanie do drugiej osoby, choć z drugiej strony posiada również powody, przez które zachowuje się bardzo ostrożnie i podejrzliwie w innych aspektach relacji. Relacja ta prawdopodobnie charakteryzuje się wielopłaszczyznową wzajemną zależnością, gdzie partnerzy mają odrębne, jak również wspólne cele. Sytuacja ta jest najpopularniejsza i najczęstsza w relacjach zawodowych, w relacjach międzyfirmowych, a także w relacjach biznesowych. Istnieją pola wspólnych interesów, lecz część informacji i zachowań jest pilnie strzeżona i niedostępna drugiej stronie. Relacje te są dojrzałe i pozostają na poziomie partnerskim.

W późniejszych latach termin „zaufanie” musiał ponownie przejść ewolucję. Jednym z czynników zmiany było coraz szersze wykorzystanie narzędzi internetowych, mediów społecznościowych, systemów rekomendacji, jak również innych systemów opartych na algorytmach uczenia maszynowego. Maszyny i komputery na tyle zdominowały codzienne działania człowieka, że zaistniała konieczność interakcji użytkownika systemów i sieci nie tylko z innymi użytkownikami, lecz także z automatami, programami, inteligentnymi agentami i rozbudowanymi systemami informatycznymi. Interakcje człowieka z maszynami stały się codziennością. Wykazujemy zaufanie do prognoz pogody przed wyjściem z domu, do informacji przekazywanych poprzez social media, do wyników otrzymanych z silników rekomendacji, które sugerują różne artykuły i usługi, opierając się na analizie ocen innych użytkowników o podobnym do naszego profilu zakupowym. Reasumując, jako społeczeństwo mamy bardzo duże zaufanie do technologii. W wielu przypadkach zaufanie użytkowników jest nieuświadomione, a większość użytkowników technologii wykorzystujących elementy sztucznej inteligencji w ogóle nie zastanawia się nad algorytmami, które determinują jej działanie. Nie mając informacji na temat roli, jaką odgrywają algorytmy sztucznej inteligencji, jesteśmy narażeni na utratę prywatności, a co gorsza podmiotowości.

Wiedza, w jaki sposób sztuczna inteligencja podejmuje swoje decyzje, jest więc kluczowa, aby możliwe były wprowadzenie poziomów zaufania i nieufności oraz ich weryfikacja. Nie mniej istotne wydają się kwestie moralne. Niezbędne jest przeanalizowanie, z punktu etycznego, podejmowania decyzji przez maszynę w kryzysowych sytuacjach. Dla przykładu jak powinna postąpić sztuczna inteligencja w sytuacji zagrożenia na drodze – czy powinna zaryzykować życie kierowcy, czy życie grupy pieszych. Problem ten jest istotnym zagadnieniem z punktu widzenia teorii i praktyki podejmowania decyzji z wykorzystaniem sztucznej inteligencji. Zgodnie z wytycznymi Komisji Europejskiej w zakresie etyki dotyczącymi godnej zaufania sztucznej inteligencji godna zaufania sztuczna inteligencja posiada trzy cechy, które muszą charakteryzować wyposażony w nią system przez cały jego cykl życia: powinna być zgodna z prawem, etyczna oraz solidna zarówno z technicznego, jak i ze społecznego punktu widzenia, ponieważ systemy SI mogą wywoływać niezamierzone szkody nawet wówczas, gdy korzysta się z nich w dobrej wierze [Komisja Europejska, 2019].

Zaufanie do sztucznej inteligencji i do systemów uczenia maszynowego zależy między innymi od transparentności użytych modeli. Modele, które nie są transparentne, z założenia powinny być wyjaśnialne i interpretowalne, aby czło-

wiek mógł zrozumieć, dlaczego otrzymuje konkretne wyniki i analizy. Ma to na celu zwiększenie transparentności tych modeli, a w konsekwencji zwiększenie zaufania użytkowników do otrzymanych wyników. W następnych punktach przedstawiono zagadnienia interpretowalności algorytmów uczenia maszynowego (ang. IML – Interpretable Machine Learning) i wytłumaczalność sztucznej inteligencji (ang. XAI – eXplainable Artificial Intelligence). W pierwszej kolejności omówiono algorytmy uczenia maszynowego w celu przedstawienia sposobu kreowania wiedzy przez maszynę i podejmowania przez nią decyzji.

## 2.2. Wiedza jako wynik uczenia maszynowego

Systemy implementujące uczenie maszynowe zapewniają możliwość automatycznego uczenia się i doskonalenia na podstawie doświadczenia (bazy uczącej) bez konieczności bezpośredniego programowania. Uczenie maszynowe koncentruje się na tworzeniu programów, które mają dostęp do danych i wykorzystują je, ucząc się samodzielnie. W efekcie możliwe jest tworzenie modeli oraz ich implementacja w odpowiedzi jedynie na określone cele działania [Duch, 1997]. Uczenie maszynowe opiera się na dziedzinach bazujących między innymi na informatyce, statystyce czy kognitywistyce. System informatyczny implementujący inteligentne algorytmy dzięki wykorzystaniu metod i algorytmów uczenia maszynowego może samodzielnie przetwarzać dane, na podstawie których tworzy i dostosowuje modele obliczeń tak, aby możliwe było zdobycie nowej wiedzy, która przełoży się na możliwość rozwiązania danego problemu. Należy jednak zaznaczyć, że nie da się przewidzieć, jakie zostaną podjęte przez nie decyzje czy wyciągnięte wnioski. Zadaniem tego typu systemów jest nauka na przykładach, co sprawia, że z upływem czasu sam autor nie jest w stanie dokładnie określić ich wiedzy.

Reprezentacja otrzymanej wiedzy poprzez implementację uczenia maszynowego zależy od samego algorytmu uczenia maszynowego. Inaczej wiedzę wykreuje algorytm drzew decyzyjnych, a inaczej algorytm sieci neuronowych. Modele uczenia maszynowego zbudowane za pomocą takich algorytmów można podzielić na dwa rodzaje: modele białej skrzynki (ang. white box) i modele czarnej skrzynki (ang. black box). Modele białej skrzynki mają strukturę, która umożliwia interpretację jego elementów. Przeglądalna struktura oraz niewielka liczba szacowanych parametrów pozwala na dokładne zrozumienie relacji wszystkich cech/zmiennych użytych w modelu. Za przykład może posłużyć

algorytm drzew decyzyjnych, dla którego wiedza jest reprezentowana za pomocą grafu zrozumiałego dla człowieka. Czytelność stworzonego grafu nie stanowi najmniejszych problemów. Przeciwnieństwem są modele czarnej skrzynki, w których struktura jest niemożliwa do interpretacji. Trudność ta wynika najczęściej ze stopnia złożoności tych modeli, liczba szacowanych parametrów oraz złożoność przekształceń matematycznych jest na tyle duża, że człowiek nie jest w stanie dostrzec wszystkich zamodelowanych przez algorytm relacji. Przykładowym algorytmem są sieci neuronowe, dla których reprezentacją wiedzy jest uzupełniona macierz wag pomiędzy neuronami. Wagi zostaną ustalone podczas etapu nauczania, proces ten jest nazywany w sieciach neuronowych retropropagacją gradientu. Niemożliwa jest jakakolwiek interpretacja wyniku, wagi wskazują na odpowiednie wyjście przy zadanym wejściu. Musimy zatem wprowadzić wytłumaczenie działania sieci neuronowej, aby wprowadzić poziomy zaufania i nieufności.

Algorytmy uczenia maszynowego nie pozyskują wiedzy w sposób analityczny, lecz korzystają z przykładów z bazy uczącej, zatem nie uzyskamy odpowiedzi na pytanie „dlaczego”. Nie poznamy procesu podjętej decyzji ani czynników, które wpłynęły na zwrotną informację. Uzyskamy wynik „ile, jak”, lecz nie będziemy znali odpowiedzi, dlaczego podany wynik jest poprawny. Statystycznie uzyskane odpowiedzi są w pełni poprawne, lecz powstaje pytanie, czy możemy im w pełni zaufać. Z tego powodu interpretowalność i wytłumaczalność tych modeli jest istotna dla uzyskania zaufania użytkowników do algorytmów uczenia maszynowego i ich wyników.

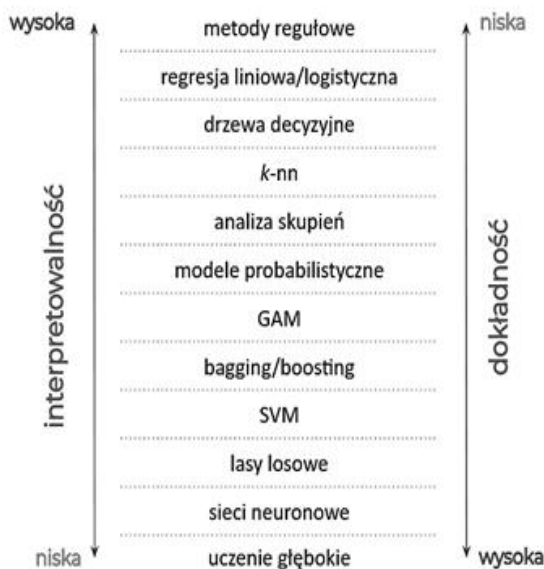
## 2.3. Wytłumaczalność modeli uczenia maszynowego

Zaufanie człowieka do sztucznej inteligencji jest ściśle związane z umiejętnością zrozumienia modeli generowanych za pomocą algorytmów maszynowego uczenia. Jak wcześniej wspomniano, część modeli poprzez swoją strukturę jest bardziej zrozumiała niż pozostałe. Brak możliwości interpretacji niektórych modeli uczenia maszynowego (brak ich transparentności) nie oznacza, że nie można wyjaśnić ich działania.

Modele typu white-box są z natury rzeczy interpretowalne i tym samym wyjaśnialne. Interpretowalność polega na opisanu mechaniki modelu, czyli zależności pomiędzy danymi a parametrami danego algorytmu. Możliwości interpretacyjne integralnie związane z konkretnym algorytmem pozwalają bada-

czom na przewidywanie zmian modelu, gdy zmienione zostaną parametry algorytmu albo dane wejściowe. Wystarczy, że ekspert dokona analizy wizualnej grafu drzewa decyzyjnego i stwierdzi, że rozumie, co się w nim dzieje, potrafi zinterpretować zachodzące relacje pomiędzy zmiennymi w modelu bez konieczności poznawania dziedziny analizowanego modelu. Interpretowalność modelu daje gwarancję rozumienia, w jaki sposób model będzie formułował prognozy, co z kolei pozwoli zrozumieć, dlaczego model popełnia błędy. Modele interpretowalne cechują się więc zrozumiałością i często komunikatywnością. Zrozumiałość jest cechą modelu mówiącą, iż człowiek jest w stanie zrozumieć sposób działania modelu bez konieczności analizy sposobu przetwarzania danych, natomiast komunikatywność to możliwość zrozumiałej dla człowieka reprezentacji wiedzy modelu [Montavon, Samek, Müller, 2018; Fernandez i in., 2019; Gleicher, 2016; Craven, 1996]. Modele typu black-box charakteryzują się niską interpretowalnością, jednak ze względu na wysoką dokładność predykcji są często stosowane w rozwiązaniach sztucznej inteligencji.

Bez względu na to, czy w procesie budowy systemów sztucznej inteligencji wykorzystuje się modele interpretowalne, czy nieinterpretowalne, bardzo ważnym zagadnieniem jest ocena jakości oszacowanych modeli. Można zauważyć, iż wraz ze wzrostem dokładności modelu możliwość jego interpretowalności znacznie maleje (zob. rys. 1) [Biecek, Burzykowski, 2021; Duval, 2019; Guo i in., 2019]. Drzewa decyzyjne ze względu na swoją konstrukcję są jednymi z najłatwiej interpretowanych i intuicyjnych algorytmów ML, niestety dokładność ich predykcji w porównaniu z innymi metodami ML jest na stosunkowo niskim poziomie. Natomiast rozwiązanie, jakim jest uczenie głębokie (ang. Deep Learning), osiąga jedne z najlepszych wyników pod względem wydajności, jednak opiera swoje działanie na podejściu black-box, przez co interpretacja działania oraz ich zrozumiałość jest najmniejsza. Ma to znacząco negatywny wpływ na zaufanie do wykorzystania tego podejścia, a w pewnych dyscyplinach ze względu na znaczną odpowiedzialność i wymagania prawne ich użycie jest niemożliwe [Rozporządzenie 2016/679; Ustawa o zmianie niektórych ustaw..., 2019, art. 43, 46, 138].



**Rys. 1.** Relacja między dokładnością modelu a jego interpretowalnością

Źródło: Opracowanie własne na podstawie: Biecek, Burzykowski [2021]; Duval [2019]; Guo i in. [2019].

Wiele technik pozwalających na ustalenie, który ze zbudowanych modeli wykazuje lepszą wytłumaczalność niezależnie od tego, czy jest interpretowalny, czy nie, nosi nazwę eXplainable Artificial Intelligence (XAI). Jedną z głównych definicji Explainable Artificial Intelligence tłumaczy je jako systemy AI, których działanie jest bardziej zrozumiałe dla człowieka poprzez dostarczanie użytkownikowi odpowiednich wyjaśnień [Gunning i in., 2019]. Koncepcja ta w głównej mierze zakłada tworzenie systemów AI, w których równie ważne, jak wydajność, jest stworzenie odpowiedniego środowiska dla użytkownika końcowego, dzięki czemu oprócz uzyskania informacji o wyniku działania algorytmu będzie on w stanie również dowiedzieć się o sposobie zastosowanego wnioskowania lub przyczynach podjęcia danej decyzji [Bellotti, Edwards, 2001]. Można więc zauważyć, że możliwość interpretacji wyników modelu interpretowalnego może być z powodzeniem wykorzystana w budowie systemu XAI jako rozszerzenie jego możliwości. Rozróżnienie między interpretacją a wytłumaczeniem (wyjaśnieniem) modeli jest istotne z punktu widzenia przygotowania strategii objaśnienia użytkownikom decyzji podjętej w ich sprawie. Aby uczynić zadość tym oczekiwaniom (a w przypadkach określonych prawnie nawet nakazom), opracowano wiele technik umożliwiających wytłumaczenie zachowania modelu.

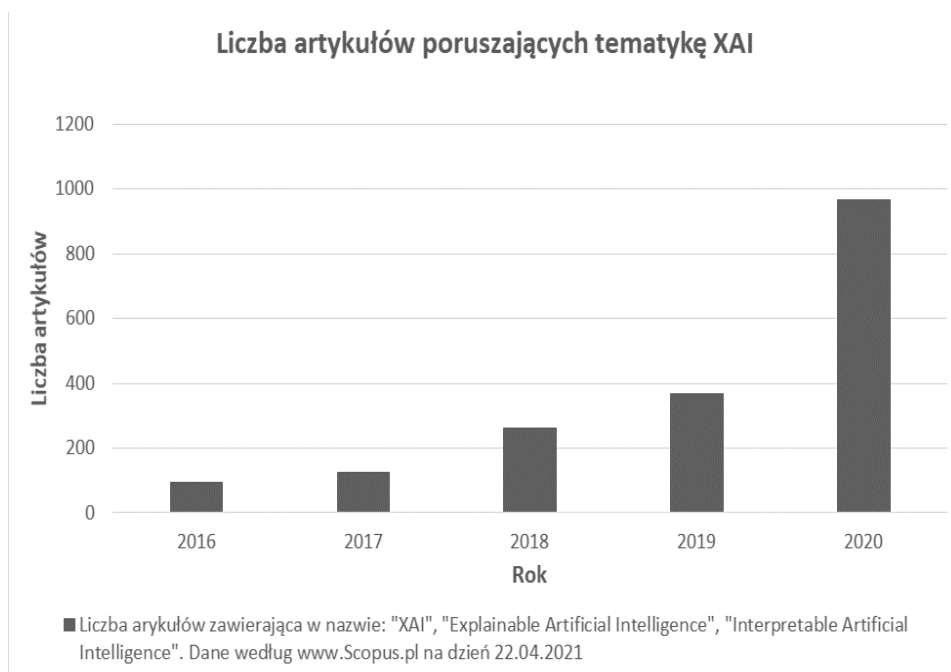
Co istotne, opracowane metody tłumaczenia zachowania modelu nieinterpretowalnego mogą zostać z powodzeniem zastosowane do modeli interpretowalnych. Dzięki temu istnieje możliwość spójnego prezentowania wyników różnych modeli oraz, co nawet ważniejsze, porównywania jakości i skuteczności wygenerowanych modeli.

Kluczowym elementem wytłumaczalności AI jest człowiek i jego miejsce w tworzeniu i wykorzystaniu AI. To jego zaufanie i zrozumienie działania algorytmów AI jest niezbędne dla ich realnego użycia w firmie. W pracy Barredo Arrieta i in. [2020] można znaleźć podział użytkowników AI ze względu na pełnione przez nich role w procesie tworzenia i wykorzystania AI:

- eksperci dziedzinowi i użytkownicy AI (np. lekarze, agenci ubezpieczeniowi) wykorzystujący wyniki działania AI; kluczowe jest ich zaufanie do modelu i zdobywanie wiedzy na podstawie wyników jego działania,
- agencje i podmioty regulacyjne oceniają modele i przygotowują regulacje prawne,
- kierownicy i osoby decyzyjne podejmują decyzje o przeznaczeniu środków na AI i opłacalności tych rozwiązań,
- badacze i twórcy AI, osoby odpowiedzialne za rozwój rozwiązań AI sprawdzają i poprawiają efektywność produktów, prowadzą badania nad nowymi możliwościami rozwoju,
- osoby dotknięte konsekwencjami decyzji opracowanych z wykorzystaniem AI.

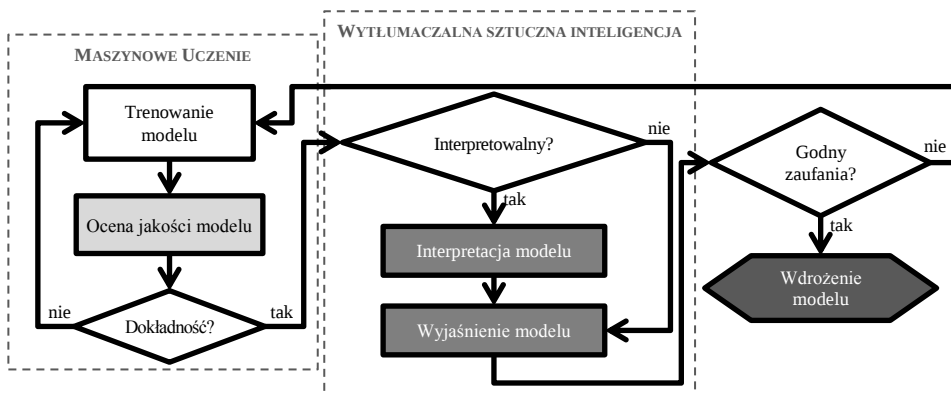
Jako główne czynniki powstrzymujące zastosowanie sztucznej inteligencji w firmie są wymieniane brak zaufania i niepewna stabilność rezultatów związana z działaniem algorytmów typu black-box. Z wymienionych powodów oraz ogromnego potencjału sztucznej inteligencji wynika rosnąca popularność tematyki związanej z XAI. Warto wziąć pod uwagę wskaźnik dotyczący liczby publikacji z zakresu wytłumaczalności sztucznej inteligencji (zob. rys. 3).

Zarówno naukowcy, jak i świat biznesu podkreślają, iż dotychczasowe podejście opierające się w głównej mierze na wydajności modeli jest niewystarczające i sygnalizują potrzebę zrozumienia dostarczanych rozwiązań. Na rysunku 2 można zaobserwować znaczący wzrost publikacji związanych z XAI. Warto również zauważyć, iż coraz więcej organizacji, tj.: AI4EU, XAI-Project, DARPA – Explainable Artificial Intelligence (XAI) Project, SoBigData, Humane AI, Google Brain czy też Explainable Machine Learning Models podejmuje działania na rzecz rozwoju XAI.



**Rys. 2.** Liczba artykułów poruszających tematykę XAI

Źródło: Opracowanie własne na podstawie [www 1].



**Rys. 3.** Wytlumaczalność sztucznej inteligencji w procesie budowy systemów SI

Źródło: Opracowanie własne.

Rysunek 3 przedstawia miejsce XAI w ścieżce implementacji algorytmów maszynowego uczenia w systemach sztucznej inteligencji ukierunkowanych na maksymalizację wartości zaufania do otrzymanego modelu. Na rysunku widać model wytrenowany przy użyciu bazy uczącej. Następnie jest realizowany etap ewaluacji modelu w celu oszacowania sprawności otrzymanego modelu. Przedstawione etapy należą do technik uczenia maszynowego i następują iteracyjnie w celu doboru najlepszego modelu z punktu widzenia jego zdolności predykcyjnych. Czasami zdarza się, że kilka modeli wykazuje podobne wyniki oceny dokładności predykcji, w takich przypadkach do kolejnego etapu można przekazać kilka najlepszych modeli, które w kolejnym etapie zostaną porównane pod kątem zaufania. Następnie rozpoczyna się faza zależna od typu algorytmu, jaki został wykorzystany w budowie modelu; jeżeli algorytm pozwala na budowę modelu interpretowalnego, następuje interpretacja oszacowanych parametrów, odkrytych wzorców i trendów. Jeżeli nie, rozpoczyna się proces wyjaśniania. Modele otrzymane w wyniku zastosowania algorytmów maszynowego uczenia przechowują bardzo rozbudowaną wiedzę. Istnieje możliwość wydobycia z nich wielu bardzo różnorodnych informacji wykorzystanych w wyjaśnianiu różnych aspektów danego modelu. Jako że głównym zadaniem modeli jest przewidywanie zachowania analizowanych obiektów, rodzą się pytania dotyczące zarówno indywidualnych osób, zjawisk, jak i całego zbioru obiektów. Ten podział ma kluczowy wpływ na proces wyjaśniania, który można uzależnić od chęci zrozumienia działania modelu w ujęciu lokalnym (skupionym na konkretnym przypadku) lub globalnym (skupionym na wielu przypadkach traktowanych jako całość). Ponadto z perspektywy zarówno lokalnej, jak i globalnej możliwe jest zrozumienie wpływu poszczególnych zmiennych wykorzystanych w budowie modelu oraz predykcji. W zależności od zastosowanej metody możliwe jest wykrycie, które ze zmiennych miały największy wpływ na ostateczny wynik predykcji lub jak zmieniłby się wynik, gdyby zmieniły się wartości tylko niektórych zmiennych (analiza *what-if*). Konsekwencją wytłumaczenia zachowania modelu może być wykrycie przyczyn niepoprawnych prognoz modelu, usunięcie z bazy wejściowej zmiennych nieistotnych lub przypadków, które wyróżniają się na tle innych, przez co sprawiają, że model daje zakłámane wyniki. Tabela 1 zawiera zestawienie najczęściej stosowanych technik w wyjaśnianiu modelu typu black-box [Biecek, Burzykowski, 2021].

**Tabela 1.** Techniki XAI

| Lokalne                                                                         | Globalne                                                                     |
|---------------------------------------------------------------------------------|------------------------------------------------------------------------------|
| <b>Czy model jest dobrze dopasowany</b>                                         |                                                                              |
| naokoło predykcji: analiza sąsiadów local-fidelity plots, local-stability plots | ogólnie: analiza reszt modelu                                                |
| <b>Jak zmienne oddziałują na</b>                                                |                                                                              |
| predykcję: Ceteris Paribus                                                      | średnią predykcję: Partial Dependence Plots, Accumulated Local Effects Plots |
| <b>Które zmienne</b>                                                            |                                                                              |
| przyczyniają się do podjęcia konkretnej decyzji: SHAP, LIME, Break Down         | są ważne dla modelu: permutacja wartości zmiennych                           |

Źródło: Opracowanie własne na podstawie: Biecek, Burzykowski [2021].

W przypadku modeli opartych na algorytmach typu black-box możliwe jest użycie jednej z dwóch kategorii metod wyjaśniania. Pierwsza grupa metod (ang. model-agnostic method) ma szerokie zastosowanie niezależne od sposobu przetwarzania danych przez model i formy prezentacji informacji. W przypadku drugiej kategorii rozwiązań (ang. model-specific method) możliwe jest ich zastosowanie jedynie na wybranych algorytmach, np. maszynie wektorów nośnych (ang. support vector machine), zespołach klasyfikatorów (ang. ensembles method), sztucznych sieciach neuronowych (ang. neural networks), konwolucyjnych sieciach neuronowych (ang. convolutional neural network) i rekurencyjnych sieciach neuronowych (ang. recurrent neural network). Metody z tej grupy posiadają odpowiednio dopasowane działanie do wymienionych algorytmów i ich struktury [Barredo Arrieta i in., 2020].

Model-agnostic method opiera swoje działanie na zastosowaniu otrzymanych wyników przez oryginalny model do budowy nowego modelu bazującego na algorytmach interpretowalnych. Takie podejście kosztem utraty części informacji uzyskuje uproszczony model, który jest dogodniejszy w zastosowaniu przez użytkownika. Jedną z najpopularniejszych metod należących do tej grupy jest LIME (ang. Local Interpretable Model-Agnostic Explanations). Metoda LIME buduje lokalne liniowe modele wokół prognozy i nie wyjaśnia całego modelu, a jedynie jego część [Ribeiro, Singh, Guastrin, 2016]. Dokonując analizy model-specific method, można zauważyć, iż są one szczególną implementacją model-agnostic method dostosowaną do wymagań (specyficzna wielowarstwowa struktura, przekształcenia wewnętrzne) konstrukcyjnych konkretnych algorytmów. Wśród sztucznych sieci neuronowych najpopularniejsze metody to: upraszczanie modelu, estymatory trafności cech, wyjaśnienia na lokalnych da-

nych. Można również spotkać opisy tekstowe oraz modele wizualne. Często stosowanym algorytmem jest Deep RED (ang. Rule Extraction from Deep Neural Networks). Polega on na ekstrakcji reguł (poziom neuronów) poprzez dodanie większej liczby drzew decyzyjnych i reguł. Ze względu na uzyskiwanie dobrych wyników konwolucyjne sieci neuronowe są często stosowane do analiz obrazu, klasyfikacji elementu i wykrywania obiektów. W ramach tych algorytmów możliwe jest wyjaśnianie na podstawie zrozumienia procesu decyzyjnego sieci oparte na porównaniu wyjść i wejść modelu oraz oceny dominujących elementów. Metodami wykorzystywanymi do tego są między innymi mapa cieplna (ang. heatmap) lub algorytm Grad-CAM (ang. Gradient-weighted Class Activation Mapping). Rozwiązania te bazują np. na wizualnej prezentacji danego obrazu z naniesionym filtrem informującym o istotności danego elementu w wyniku. Możliwa jest również analiza bazująca na podziale poszczególnych warstw odpowiedzialnych za rozpoznawanie poszczególnych elementów. Końcowy wynik pozwala na identyfikację, która z warstw jest odpowiedzialna za rozpoznawanie elementów z daną szczegółowością. Rekurencyjne sieci neuronowe stosowane w analizie danych sekwencyjnych, przetwarzaniu języka naturalnego i analizie szeregów czasowych mogą być wyjaśniane z użyciem metod istotności cech lub lokalnej eksploracji danych [Barredo Arrieta i in., 2020].

Analizując sposoby wyjaśniania modeli na przykładzie metody LIME, można zauważyć, iż pozwala ona na budowę nowego modelu na podstawie wyników modelu pierwotnego. Opracowany model jest wtórny, a jego wyniki są przybliżone do oryginalnego (część informacji jest tracona), jednak jest możliwa jego interpretacja. Budowa modelu wyjaśniającego polega na wyznaczeniu funkcji regresji liniowej na podstawie losowo wybranych przypadków znajdujących się w pobliżu (lokalnie) przypadku prognozowanego. Parametry tej funkcji pozwalają na wskazanie znaczenia cech, na podstawie których podejmowana jest decyzja modelu typu black-box [Ribeiro, Singh, Guestrin, 2016]. Przykładowa analiza z zastosowaniem LIME została zaprezentowana w pracy Biecek i Burzykowski [2021]. Dla danych dotyczących brytyjskiego okrętu RMS Titanic, który zderzył się z górą lodową i zatonął, przeprowadzono klasyfikację przeżycia potencjalnego pasażera. Opracowane modele wykorzystywały dane osób biorących udział w rejsie związane z wiekiem, płcią czy klasą biletu. Dzięki zastosowaniu LIME można zaobserwować, iż wiek poniżej 16 lat w blisko 24% wpłynął negatywnie na szansę przeżycia, natomiast płeć (29%) i klasa biletu (44%) w znacznym stopniu zwiększały szansę na przeżycie pasażera. W przypadku modeli black-box po podaniu parametrów uzyskuje się informację, czy

dana osoba zostanie zaklasyfikowana do osób, które przeżyły, czy też nie, jednak nie dostarczają one informacji, która bądź które z cech miały na to wpływ i w jakim stopniu. Po zastosowaniu odpowiedniego algorytmu, w tym przypadku LIME, możliwe jest uzyskanie przybliżonego obrazu zjawiska i ocena wpływu poszczególnych cech na wynik klasyfikacji.

W trakcie przygotowywania modelu uczenia maszynowego pod uwagę powinno się brać zarówno miary oceny jakości poszczególnych modeli, jak i możliwość interpretowalności oraz wytłumaczalności poszczególnych algorytmów i ich wyników. Takie podejście pozwala na lepsze zrozumienie procesu, który został zamodelowany, a jest to możliwe z zastosowaniem zaprezentowanych technik XAI. Zastosowanie takiego rozwiązania pozwala na uzyskanie większej wiedzy zarówno przez odbiorców końcowych danego modelu, jak i innych uczestników całego procesu.

## Podsumowanie

Znaczna popularność rozwiązań dotyczących sztucznej inteligencji, w tym uczenia maszynowego, powoduje rosnące zainteresowanie i stosowanie tych metod. Wykorzystanie tego potencjału jest osiągalne z zastosowaniem odpowiednich technik wytłumaczalności, dzięki czemu możliwe jest pogłębienie wiedzy i użycie w szerszym zakresie przypadków. Aby możliwe było zaufanie metodom uczenia maszynowego, muszą się one cechować zgodnością z prawem, etycznością oraz rzetelnością.

W niniejszym rozdziale przedstawiono rolę poziomu zaufania do algorytmów uczenia maszynowego, której ważnym elementem jest ich wytłumaczalność. Do pełnego zastosowania wiedzy, jaka drzemie w sztucznej inteligencji, wymagane jest jej odpowiednie zrozumienie i poznanie elementów, na których opierają swoje działanie poszczególne algorytmy. Takie podejście pozwala na zastosowanie sztucznej inteligencji w coraz szerszym zakresie. Wraz ze wzrostem zrozumienia stosowanych metod rośnie zaufanie do nich. Zaufanie do systemów jest również zależne od ich wytłumaczalności i interpretowalności.

Podejście eXplainable Artificial Intelligence pozwala na tworzenie systemów AI, w których nie tylko wydajność algorytmu uczenia maszynowego czy precyzja modelu w podejmowaniu decyzji stanowią obszary troski, ale ważne jest także stworzenie odpowiedniego środowiska dla użytkownika końcowego. Jak wskazano w rozdziale, złożoność modeli będących wynikiem zastosowania

różnorodnych algorytmów maszynowego uczenia jest tak duża, że konieczne staje się opracowanie nowych metod objaśniania ich wyników oraz prezentowania sposobu, w jaki model podejmuje decyzje.

Stosowane modele uczenia maszynowego oferują coraz większą dokładność, jednak w większości przypadków powoduje to zmniejszenie możliwości interpretowalności przyczyn, na podstawie których został podjęty dany wniosek ze względu na sposób działania i skomplikowanie algorytmu. W zależności od typu wykorzystywanych algorytmów dostępne są różne metody pozwalające na wyjaśnianie modeli model-agnostic method, model-specific method. Metody te są dostosowane do konkretnych algorytmów lub ich grup i pozwalają między innymi na budowę nowego uproszczonego, ale wyjaśnialnego modelu z wykorzystaniem pierwotnie uzyskanego rozwiązania. Przykładami tego typu rozwiązań mogą być metody LIME i SHAPE.

Zarówno twórcy, jak i użytkownicy sztucznej inteligencji stawiają nowe wyzwania przed algorytmami, a dzięki XAI algorytmy stają się już nie tylko efektywne, ale również bardziej zrozumiałe i interpretowalne dla człowieka, dzięki czemu może on im zaufać i czerpać z nich nową wiedzę. Należy jednak pamiętać, iż mimo znaczącego wzrostu zainteresowania tematyką XAI nadal można wymienić wiele wyzwań i trudności stojących przed badaczami, do których można zaliczyć na przykład wybór pomiędzy koncentracją na człowieku a maszynie, położenie nacisku na dokładność lub interpretowalność czy też wyjaśnianie zdolności lub decyzji AI [Gunning i in., 2019].

## Literatura

- Barber B. (1983), *The Logic and Limits of Trust*, Rutgers University Press, New Brunswick, NJ.
- Barredo Arrieta A., Díaz-Rodríguez N., Del Ser J., Bennetot A., Tabik S., Barbado A., García S., Gil-Lopez S., Molina D., Benjamins R., Chatila R., Herrera F. (2020), *Explainable Artificial Intelligence (XAI): Concepts, Taxonomies, Opportunities and Challenges toward Responsible AI*, "Information Fusion", Vol. 58, s. 82-115, <https://doi.org/https://doi.org/10.1016/j.in us.2019.12.012>.
- Bellotti V., Edwards K. (2001), *Intelligibility and Accountability: Human Considerations in Context-aware Systems*, "Human Computer Interaction", Vol. 16(2-4), s. 193-212, [https://doi.org/10.1207/S15327051HCI16234\\_05](https://doi.org/10.1207/S15327051HCI16234_05).
- Biecek P., Burzykowski T. (2021), *Explanatory Model Analysis: Explore, Explain and Examine Predictive Models*, Chapman & Hall/CRC, New York.

- Craven M.W. (1996), *Extracting Comprehensible Models from Trained Neural Networks (tech. rep.)*, University of Wisconsin-Madison Department of Computer Sciences, <https://minds.wisconsin.edu/bitstream/handle/1793/60078/TR1326.pdf?sequence=1>.
- Duch W. (1997), *Fascynujący świat komputerów*, Wydawnictwo Nakom, Poznań.
- Duval A. (2019), *Explainable Artificial Intelligence (XAI)*, MA4K9 Scholarly Report, The University of Warwick, Mathematics Institute.
- Fernandez A., Herrera F., Cordon O., Jose del Jesus M., Marcelloni F. (2019), *Evolutionary Fuzzy Systems for Explainable Artificial Intelligence: Why, When, What for, and Where to?* "IEEE Computational Intelligence Magazine", Vol. 14(1), s. 69-81.
- Gall R. (2018), *Machine Learning Explainability vs Interpretability: Two Concepts that Could Help Restore Trust in AI*, <https://www.Kdnuggets.com/2018/12/machine-learning-explainability-interpretability-ai.html>.
- Gleicher M. (2016), *A Framework for Considering Comprehensibility in Modeling* [PMID: 27441712], "Big Data", Vol. 4(2), s. 75-88, <https://doi.org/10.1089/big.2016.0007>.
- Goodman B., Flaxman S. (2017), *European Union Regulations on Algorithmic Decision-making and a Right to Explanation*, "AI Magazine", Vol. 38(3), s. 50-57, <https://doi.org/10.1609/aimag.v38i3.2741>.
- Gunning D., Stefik M., Choi J., Miller T., Stumpf S., Yang G.-Z. (2019), *XAI-Explainable Artificial Intelligence*, "Science Robotics", Vol. 4(37), <https://doi.org/10.1126/scirobotics.aay7120>.
- Guo M., Zhang Q., Liao X., Chen Y. (2019), *An Interpretable Machine Learning Framework for Modelling Human Decision Behavior*, arXiv preprint arXiv:1906.01233.
- Komisja Europejska (2019), *Wytyczne w zakresie etyki dotyczące godnej zaufania sztucznej inteligencji*, [https://www.europarl.europa.eu/meetdocs/2014\\_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI\\_PL.pdf](https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI_PL.pdf).
- Lewicki R.J., McAllister D.J., Bies R.J. (1998), *Trust and Distrust: New Relationships and Realities*, "Academy of Management Review", Vol. 23(3), s. 438-458.
- Mayer R.C., Davis J.H., Schoorman F.D. (1995), *An Integrative Model of Organizational Trust*, "Academy of Management Review", Vol. 20(3), s. 709-734.
- Mellinger G.D. (1956), *Interpersonal Trust as a Factor in Communication*, "The Journal of Abnormal and Social Psychology", Vol. 52(3), s. 304-309.
- Molnar C. (2020), *Interpretable Machine Learning*, [www.lulu.com](http://www.lulu.com) (dostęp: 23.04.2021).
- Montavon G., Samek W., Müller K.-R. (2018), *Methods for Interpreting and Understanding Deep Neural Networks*, "Digital Signal Processing", Vol. 73, s. 1-15, <https://doi.org/https://doi.org/10.1016/j.dsp.2017.10.011>.
- Read W.H. (1962), *Upward Communication in Industrial Hierarchies*, "Human Relations", Vol. 15(1), s. 3-15.

- Ribeiro M.T., Singh S., Guestrin C. (2016), “*Why Should I Trust You?*”: *Explaining the Predictions of any Classifier*, Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, s. 1135-1144, <https://doi.org/10.1145/2939672.2939778>.
- Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2016/679 z dnia 27 kwietnia 2016 roku w sprawie ochrony osób fizycznych w związku z przetwarzaniem danych osobowych i w sprawie swobodnego przepływu takich danych oraz uchylenia dyrektywy 95/46/WE. Dziennik Urzędowy UE L119, s. 1-88, Baza Aktów Prawnych Unii Europejskiej [online], <http://eur-lex.europa.eu> (dostęp: 23.04.2021).
- Siau K., Wang W. (2018), *Building Trust in Artificial Intelligence, Machine Learning, and Robotics*, “Cutter Business Technology Journal”, Vol. 31(2), s. 47-53.
- Toreini E., Aitken M., Coopamootoo K., Elliott K., Zelaya C.G., van Moorsel A. (2020), *The Relationship between Trust in AI and Trustworthy Machine Learning Technologies*, Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency, s. 272-283.
- Ustawa z dnia 21 lutego 2019 roku o zmianie niektórych ustaw w związku z zapewnieniem stosowania rozporządzenia Parlamentu Europejskiego i Rady (UE) 2016/679 z dnia 27 kwietnia 2016 roku w sprawie ochrony osób fizycznych w związku z przetwarzaniem danych osobowych i w sprawie swobodnego przepływu takich danych oraz uchylenia dyrektywy 95/46/WE (ogólne rozporządzenie o ochronie danych). Dz.U. 2019, poz. 730.
- [www 1] [www.scopus.pl](http://www.scopus.pl) (dostęp: 22.04.2021).





WYDAWNICTWO UNIWERSYTETU EKONOMICZNEGO W KATOWICACH

---

NOWOŚCI WYDAWNICZE

*Doskonalenie wykorzystywania danych w organizacjach non profit.* Joanna Palonka. Monografia, 206 s., ISBN 978-83-7875-717-7.

*Internacjonalne relacje we współczesnym świecie.* Maria Balcerowicz-Szkutnik, Anna Sączewska-Piotrowska (red.). Monografia, 190 s., ISBN 978-83-7875-711-5.

*Kraje Grupy Wyszehradzkiej i ich gospodarka a praca, aktywność zawodowa i przedsiębiorczość młodego pokolenia. Analiza porównawcza.* Urszula Swadźba, Stanisław Swadźba. Monografia, 222 s., ISBN 978-83-7875-704-7.

*Media w procesie tworzenia i realizacji polityki regionalnej w województwie śląskim.* Zbigniew Widera (red.). Monografia, 162 s., ISBN 978-83-7875-719-1.

*Modeling of complex data sets and risk analysis.* Grażyna Trzpiot (ed.). Monografia, 148 s., ISBN 978-83-7875-701-6.

*Modelowanie nieostre w wielokryterialnym wspomaganiu decyzji inwestycyjnych.* Ewa Pośpiech. Monografia, 106 s., ISBN 978-83-7875-691-0.

*Nowe media i technologie w komunikacji marketingowej – wybrane obszary zastosowań.* Andrzej Bajdak, Zbigniew Spyra (red.). Monografia, 104 s., ISBN 978-83-7875-693-4.

*Praktyka zarządzania wiedzą klienta w perspektywie małych i średnich przedsiębiorstw usługowych.* Agnieszka Dziubińska, Marcin Komańda. Monografia, 110 s., ISBN 978-83-7875-688-0.

*Trendy w badaniu przedsiębiorczości.* Katarzyna Bratnicka-Mysłiwiec, Monika Kulikowska-Pawlak (red.). Monografia, 126 s., ISBN 978-83-7875-709-2.

*Trzydzieści lat transformacji.* Jacek Pietrucha, Jolanta Pasioneck (red.). Monografia, 194 s., ISBN 978-83-7875-706-1.

*Zachowania strategiczne przedsiębiorstw w warunkach rynków wylaniających się.* Agnieszka Dziubińska. Monografia, 160 s., ISBN 978-83-7875-696-5.

*Zasobowe uwarunkowania podejmowania decyzji w organizacjach publicznych na przykładzie szpitali.* Karolina Szymaniec-Mlicka. Monografia, 124 s., ISBN 978-83-7875-699-6.

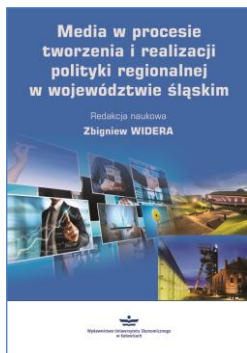
*Zastosowanie analizy klas ukrytych w badaniach ekonomicznych.* Justyna Brzezińska. Monografia, 128 s., ISBN 978-83-7875-713-9.

*Złota stal. Restrukturyzacja i prywatyzacja polskiego hutnictwa żelaza i stali.* Jerzy Podsiadło. Monografia, 230 s., ISBN 978-83-7875-686-6.



## WYDAWNICTWO UNIWERSYTETU EKONOMICZNEGO W KATOWICACH

### NOWOŚCI WYDAWNICZE



*Media w procesie tworzenia i realizacji polityki regionalnej w województwie śląskim*  
Zbigniew Widera (red.)

Publikacja zawiera wieloaspektową refleksję nad rolą mediów regionalnych w kształtowaniu właściwych postaw obywatelskich, będących determinantą tworzenia i realizacji polityki regionalnej w województwie śląskim. Prezentowana pozycja wydawnicza jest efektem projektu badawczego realizowanego w Kolegium Informatyki i Komunikacji Uniwersytetu Ekonomicznego w Katowicach, skoncentrowanego na roli środków masowego przekazu w kształtowaniu postaw obywateli, przyczyniających się do właściwie realizowanej polityki regionalnej. Rezultaty rozważań o charakterze interdyscyplinarnym mogą być przydatne zarówno w edukacji, jak i w zarządzaniu województwem śląskim na różnych jego szczeblach. Monografia, 162 s., ISBN 978-83-7875-719-1.



*Trzydzieści lat transformacji*  
Jacek Pietrucha, Jolanta Pasioneck (red.)

Publikacja prezentuje rozważania dotyczące przebiegu i skutków transformacji systemowej w Polsce i innych krajach Europy Środkowej i Wschodniej. Przedmiotem zainteresowania Autorów stał się zwłaszcza ład instytucjonalny, zarówno jego docelowa postać, jak i doświadczenia w zakresie jego reformowania, a także szczegółowe aspekty transformacji systemowej w wybranych obszarach: innowacje i badania naukowe, kapitał ludzki, produktywność pracy, rozwój społecznie zrównoważony czy system finansowy. Równolegle do transformacji systemowej miały miejsce znaczące zmiany dotyczące charakteru nauk ekonomicznych oraz gospodarki światowej, stanowiące istotny punkt odniesienia dla zmian następujących w Europie Środkowej i Wschodniej, czemu również przyjrano się w niniejszej monografii. Monografia, 194 s., ISBN 978-83-7875-706-1.



*Kraje Grupy Wyszehradzkiej i ich gospodarka a praca, aktywność zawodowa i przedsiębiorczość młodego pokolenia. Analiza porównawcza*  
Urszula Swadźba, Stanisław Swadźba

Książka napisana przez Autorów dwóch specjalności – ekonomistę i socjologę – prezentuje i omawia społeczeństwo i gospodarkę krajów Grupy Wyszehradzkiej ze szczególnym uwzględnieniem problematyki młodego pokolenia. Ukazano w niej stan świadomości ekonomicznej członków V4 w aspekcie pracy zawodowej, aktywności gospodarczej i przedsiębiorczości, pokazano pozytywne i negatywne skutki podejmowania pracy zawodowej przez studentów i ich motywacje ku temu, jak również zanalizowano plany zawodowe tej grupy społecznej i jej mobilność. W rezultacie rozważań przedstawiono podobieństwa i różnice występujące między młodzieżą krajów Grupy Wyszehradzkiej. Publikacja powstała we współpracy z Uniwersytetem Śląskim w Katowicach. Monografia, 222 s., ISBN 978-83-7875-704-7.